

MICROCHIP FABRICATION

A PRACTICAL GUIDE TO SEMICONDUCTOR PROCESSING

Peter Van Zant | **Sixth Edition**

VISIT...

LANZAROTE
Caliente.COM

Microchip Fabrication

About the Author

Peter Van Zant has a long resume from the semiconductor industry. He started in an IBM research facility in New York State and worked his way to Silicon Valley by way of Texas Instruments in Dallas, Texas. In Silicon Valley, he held wafer-fabrication process engineering and management positions, including National Semiconductor and Monolithic Memories. He was an instructor at Foothill College, Los Altos, California, where he taught introductory semiconductor classes and advanced programs for starting process engineers. Peter and his wife Mary DeWitt founded Semiconductor Services, a wafer fabrication-oriented training and consulting company and then wrote and published the first edition of *Microchip Fabrication*. McGraw-Hill has published all of the subsequent editions. Peter and Mary sold Semiconductor Services and moved to the Sierra Nevada foothills where he founded Peter Van Zant Associates, providing consulting services to industry and the legal profession. Peter also served two terms as a Nevada County Supervisor and currently provides consulting to Sierra conservation organizations.

Microchip Fabrication

Peter Van Zant

Sixth Edition



New York Chicago San Francisco
Athens London Madrid
Mexico City Milan New Delhi
Singapore Sydney Toronto

Copyright © 2014 by McGraw-Hill Education. All rights reserved. Except as permitted under the United States Copyright Act of 1976, no part of this publication may be reproduced or distributed in any form or by any means, or stored in a database or retrieval system, without the prior written permission of the publisher.

ISBN: 978-0-07-182102-5

MHID: 0-07-182102-3

e-Book conversion by Cenveo® Publisher Services

Version 1.0

The material in this eBook also appears in the print version of this title: ISBN: 978-0-07-182101-8, MHID: 0-07-182101-5.

McGraw-Hill Education eBooks are available at special quantity discounts to use as premiums and sales promotions, or for use in corporate training programs. To contact a representative, please visit the Contact Us page at www.mhprofessional.com.

All trademarks are trademarks of their respective owners. Rather than put a trademark symbol after every occurrence of a trademarked name, we use names in an editorial fashion only, and to the benefit of the trademark owner, with no intention of infringement of the trademark. Where such designations appear in this book, they have been printed with initial caps.

Information has been obtained by McGraw-Hill Education from sources believed to be reliable. However, because of the possibility of human or mechanical error by our sources, McGraw-Hill Education, or others, McGraw-Hill Education does not guarantee the accuracy, adequacy, or completeness of any information and is not responsible for any errors or omissions or the results obtained from the use of such information.

TERMS OF USE

This is a copyrighted work and McGraw-Hill Education and its licensors reserve all rights in and to the work. Use of this work is subject to these terms. Except as permitted under the Copyright Act of 1976 and the right to store and retrieve one copy of the work, you may not decompile, disassemble, reverse engineer, reproduce, modify, create derivative works based upon, transmit, distribute, disseminate, sell, publish or sublicense the work or any part of it without McGraw-Hill Education's prior consent. You may use the work for your own noncommercial and personal use; any other use of the work is strictly prohibited. Your right to use the work may be terminated if you fail to comply with these terms.

THE WORK IS PROVIDED "AS IS." McGRAW-HILL EDUCATION AND ITS LICENSORS MAKE NO GUARANTEES OR WARRANTIES AS TO THE ACCURACY, ADEQUACY OR COMPLETENESS OF OR RESULTS TO BE OBTAINED FROM USING THE WORK, INCLUDING ANY INFORMATION THAT CAN BE ACCESSED THROUGH THE WORK VIA HYPERLINK OR OTHERWISE, AND EXPRESSLY DISCLAIM ANY WARRANTY, EXPRESS OR IMPLIED, INCLUDING BUT NOT LIMITED TO IMPLIED WARRANTIES OF MERCHANTABILITY OR FITNESS FOR A PARTICULAR PURPOSE. McGraw-Hill Education and its licensors do not warrant or guarantee that the functions contained in the work will meet your requirements or that its operation will be uninterrupted or error free. Neither McGraw-Hill Education nor its licensors shall be liable to you or anyone else for any inaccuracy, error or omission, regardless of cause, in the work or for any damages resulting therefrom. McGraw-Hill Education has no responsibility for the content of any information accessed through the work. Under no circumstances shall McGraw-Hill Education and/or its licensors be liable for any indirect, incidental, special, punitive, consequential or similar damages that result from the use of or inability to use the work, even if any of them has been advised of the possibility of such damages. This limitation of liability shall apply to any claim or cause whatsoever whether such claim or cause arises in contract, tort or otherwise.

The Sixth Edition of *Microchip Fabrication* is dedicated to my sons and their families: son Patrick and wife Cindy King and granddaughter Rebecca; son Jeffrey, my hiking and adventuring buddy; son Stephen and wife Antionette McKinnley, and granddaughter Kristina and grandson Kyle.

They have been part of my semiconductor odyssey from my career start at IBM in New York State, through Texas at Texas Instruments, to settling into the high-tech hub of Silicon Valley. They have shared the rewards of this fabled industry and the absences that were due to long shifts, weekends, and travel to the far-flung centers of the worldwide semiconductor industry.

Thank you, and I love you all!

Contents

[Preface](#)

[**1 The Semiconductor Industry**](#)

[Introduction](#)

[Birth of an Industry](#)

[The Solid-State Era](#)

[Integrated Circuits \(ICs\)](#)

[Process and Product Trends](#)

[Moore's Law](#)

[Decreasing Feature Size](#)

[Increasing Chip and Wafer Size](#)

[Reduction in Defect Density](#)

[Increase in Interconnection Levels](#)

[The Semiconductor Industry Association Roadmap](#)

[Chip Cost](#)

[Industry Organization](#)

[Stages of Manufacturing](#)

[Six Decades of Advances in Microchip Fabrication Processes](#)

[The Nano Era](#)

[Review Topics](#)

[References](#)

[**2 Properties of Semiconductor Materials and Chemicals**](#)

[Introduction](#)

[Atomic Structure](#)

[The Bohr Atom](#)

[The Periodic Table of the Elements](#)

[Electrical Conduction](#)

[Conductors](#)

[Dielectrics and Capacitors](#)

[Resistors](#)

[Intrinsic Semiconductors](#)

[Doped Semiconductors](#)

[Electron and Hole Conduction](#)

[Carrier Mobility](#)

Semiconductor Production Materials

Germanium and Silicon

Semiconducting Compounds

Silicon Germanium

Engineered Substrates

Ferroelectric Materials

Diamond Semiconductors

Process Chemicals

Molecules, Compounds, and Mixtures

Ions

States of Matter

Solids, Liquids, and Gases

Plasma State

Properties of Matter

Temperature

Density, Specific Gravity, and Vapor Density

Pressure and Vacuum

Acids, Alkalis, and Solvents

Acids and Alkalis

Solvents

Chemical Purity and Cleanliness

Safety Issues

The Material Safety Data Sheet

Review Topics

References

3 Crystal Growth and Silicon Wafer Preparation

Introduction

Semiconductor Silicon Preparation

Silicon Wafer Preparation Stages

Crystalline Materials

Unit Cells

Poly and Single Crystals

Crystal Orientation

Crystal Growth

Czochralski Method

Liquid-Encapsulated Czochralski

Float Zone

Crystal and Wafer Quality

Point Defects

Dislocations

- [Growth Defects](#)
- [Wafer Preparation](#)
 - [End Cropping](#)
 - [Diameter Grinding](#)
 - [Crystal Orientation, Conductivity, and Resistivity Check](#)
 - [Grinding Orientation Indicators](#)
- [Wafer Slicing](#)
- [Wafer Marking](#)
- [Rough Polish](#)
- [Chemical Mechanical Polishing](#)
- [Backside Processing](#)
- [Double-Sided Polishing](#)
- [Edge Grinding and Polishing](#)
- [Wafer Evaluation](#)
- [Oxidation](#)
- [Packaging](#)
 - [Wafer Types and Uses](#)
 - [Reclaim Wafers](#)
- [Engineered Wafers \(Substrates\)](#)
- [Review Topics](#)
- [References](#)

[4 Overview of Wafer Fabrication and Packaging](#)

- [Introduction](#)
- [Goal of Wafer Fabrication](#)
- [Wafer Terminology](#)
- [Chip Terminology](#)
- [Basic Wafer-Fabrication Operations](#)
- [Layering](#)
 - [Patterning](#)
 - [Circuit Design](#)
 - [Reticle and Masks](#)
 - [Doping](#)
 - [Heat Treatments](#)
- [Example Fabrication Process](#)
- [Wafer Sort](#)
- [Packaging](#)
- [Summary](#)
- [Review Topics](#)
- [References](#)

5 Contamination Control

Introduction

The Problem

Contamination-Caused Problems

Contamination Sources

General Sources

Air

Clean Air Strategies

Cleanroom Workstation Strategy

Tunnel or Bay Concept

Micro-and Mini-Environments

Temperature, Humidity, and Smog

Cleanroom Construction

Construction Materials

Cleanroom Elements

Personnel-Generated Contamination

Process Water

Process Chemicals

Equipment

Cleanroom Materials and Supplies

Cleanroom Maintenance

Wafer-Surface Cleaning

Particulate Removal

Wafer Scrubbers

High-Pressure Water Cleaning

Organic Residues

Inorganic Residues

Chemical-Cleaning Solutions

General Chemical Cleaning

Oxide Layer Removal

Room Temperature and Ozonated Chemistries

Water Rinsing

Drying Techniques

Contamination Detection

Review Topics

References

6 Productivity and Process Yields

Overview

Yield Measurement Points

Accumulative Wafer-Fabrication Yield

Wafer-Fabrication Yield Limiters

Number of Process Steps

Wafer Breakage and Warping

Process Variation

Mask Defects

Wafer-Sort Yield Factors

Wafer Diameter and Edge Die

Wafer Diameter and Die Size

Wafer Diameter and Crystal Defects

Wafer Diameter and Process Variations

Die Area and Defect Density

Circuit Density and Defect Density

Number of Process Steps

Feature Size and Defect Size

Process Cycle Time

Wafer-Sort Yield Formulas

Assembly and Final Test Yields

Overall Process Yields

Review Topics

References

7 Oxidation

Introduction

Silicon Dioxide Layer Uses

Surface Passivation

Doping Barrier

Surface Dielectric

Device Dielectric (MOS Gates)

Device Oxide Thicknesses

Thermal Oxidation Mechanisms

Influences on the Oxidation Rate

Thermal Oxidation Methods

Horizontal Tube Furnaces

Temperature Control System

Source Cabinet

Vertical Tube Furnaces

Rapid Thermal Processing

High-Pressure Oxidation

Oxidant Sources

Oxidation Processes

Preoxidation Wafer Cleaning

[Postoxidation Evaluation](#)

[Surface Inspection](#)

[Oxide Thickness](#)

[Oxide and Furnace Cleanliness](#)

[Thermal Nitridation](#)

[Review Topics](#)

[References](#)

[8 The Ten-Step Patterning Process—Surface Preparation to Exposure](#)

[Introduction](#)

[Overview of the Photomasking Process](#)

[Ten-Step Process](#)

[Basic Photoresist Chemistry](#)

[Photoresist](#)

[Photoresist Performance Factors](#)

[Resolution Capability](#)

[Adhesion Capability](#)

[Process Latitude](#)

[Pinholes](#)

[Particle and Contamination Levels](#)

[Step Coverage](#)

[Thermal Flow](#)

[Comparison of Positive and Negative Resists](#)

[Physical Properties of Photoresists](#)

[Solids Content](#)

[Viscosity](#)

[Surface Tension](#)

[Index of Refraction](#)

[Storage and Control of Photoresists](#)

[Light and Heat Sensitivity](#)

[Viscosity Sensitivity](#)

[Shelf Life](#)

[Cleanliness](#)

[Photomasking Processes—Surface Preparation to Exposure](#)

[Surface Preparation](#)

[Particle Removal](#)

[Dehydration Baking](#)

[Wafer Priming](#)

[Spin Priming](#)

[Vapor Priming](#)

[Photoresist Application \(Spinning\)](#)

[The Static Dispense Spin Process](#)

[Dynamic Dispense](#)

[Moving-Arm Dispensing](#)

[Manual Spinners](#)

[Automatic Spinners](#)

[Edge Bead Removal](#)

[Backside Coating](#)

[Soft Bake](#)

[Convection Ovens](#)

[Manual Hot Plates](#)

[In-Line, Single-Wafer Hot Plates](#)

[Moving-Belt Hot Plates](#)

[Moving-Belt Infrared Ovens](#)

[Microwave Baking](#)

[Vacuum Baking](#)

[Alignment and Exposure](#)

[Alignment and Exposure Systems](#)

[Exposure Sources](#)

[Alignment Criteria](#)

[Aligner Types](#)

[Postexposure Bake](#)

[Advanced Lithography](#)

[Review Topics](#)

[References](#)

[9 The Ten-Step Patterning Process—Developing to Final Inspection](#)

[Introduction](#)

[Development](#)

[Positive Resist Development](#)

[Negative Resist Development](#)

[Wet Development Processes](#)

[Dry \(or Plasma\) Development](#)

[Hard Bake](#)

[Hard-Bake Methods](#)

[Hard-Bake Process](#)

[Develop Inspect](#)

[Develop Inspect Reject Categories](#)

[Develop Inspect Methods](#)

[Causes for Rejecting at the Develop Inspection Stage](#)

[Etch](#)

[Wet Etching](#)

- [Etch Goals and Issues](#)
- [Incomplete Etch](#)
- [Overetch and Undercutting](#)
- [Selectivity](#)
- [Wet-Spray Etching](#)
- [Silicon Wet Etching](#)
- [Silicon Dioxide Wet Etching](#)
- [Aluminum-Film Wet Etching](#)
- [Deposited-Oxide Wet Etching](#)
- [Silicon Nitride Wet Etching](#)
- [Vapor Etching](#)

[Dry Etch](#)

- [Plasma Etching](#)
- [Etch Rate](#)
- [Radiation Damage](#)
- [Selectivity](#)
- [Ion-Beam Etching](#)
- [Reactive Ion Etching](#)

[Resist Effects in Dry Etching](#)

[Resist Stripping](#)

- [Wet Chemical Stripping of Nonmetallized Surfaces](#)
- [Wet Chemical Stripping of Metallized Surfaces](#)
- [Dry Stripping](#)
- [Post-Ion Implant and Plasma Etch Stripping](#)

[New Stripping Challenges](#)

[Final Inspection](#)

[Mask Making](#)

[Summary](#)

[Review Topics](#)

[References](#)

[10 Next Generation Lithography](#)

[Introduction](#)

[Challenges of Next Generation Lithography](#)

- [High-Pressure Mercury Lamp Sources](#)
- [Excimer Lasers](#)
- [Extreme Ultraviolet](#)
- [X-Rays](#)
- [Electron Beam or Direct Writing](#)
- [Numerical Aperture of a Lens](#)

[Other Exposure Issues](#)

- [Variable Numerical Aperture Lenses](#)
- [Immersion Exposure System](#)
- [Amplified Resist](#)
- [Contrast Effects](#)
- [Other Resolution Challenges and Solutions](#)
 - [Off-Axis Illumination](#)
 - [Lens Issues and Reflection Systems](#)
 - [Phase-Shift Masks](#)
 - [Optical Proximity Corrected or Optical Process Correction](#)
 - [Annular-Ring Illumination](#)
 - [Pellicles](#)
- [Surface Problems](#)
 - [Resist Light Scattering](#)
 - [Subsurface Reflectivity](#)
- [Antireflective Coatings](#)
 - [Standing Waves](#)
 - [Planarization](#)
- [Photoresist Process Advances](#)
 - [Multilayer Resist or Surface Imaging](#)
 - [Silylation or DESIRE Process](#)
 - [Polyimide Planarization Layers](#)
 - [Etchback Planarization](#)
 - [Dual-Damascene Process](#)
 - [Chemical Mechanical Polishing](#)
 - [Slurry](#)
 - [Polishing Rates](#)
 - [Planarity](#)
 - [Post-CMP Clean](#)
 - [CMP Tools](#)
 - [CMP Summary](#)
 - [Reflow](#)
 - [Image Reversal](#)
 - [Contrast Enhancement Layers](#)
 - [Dyed Resists](#)
- [Improving Etch Definition](#)
 - [Lift-Off Process](#)
- [Self-Aligned Structures](#)
- [Etch Profile Control](#)
- [Review Topics](#)
- [References](#)

11 Doping

[Introduction](#)

[The Diffusion Concept](#)

[Formation of a Doped Region and Junction](#)

[The N-P Junction](#)

[Doping Process Goals](#)

[Graphical Representation of Junctions](#)

[Concentration versus Depth Graphs](#)

[Lateral Diffusion](#)

[Same-Type Doping](#)

[Diffusion Process Steps](#)

[Deposition](#)

[Lateral Diffusion](#)

[Same-Type Doping](#)

[Dopant Sources](#)

[Drive-In Oxidation](#)

[Oxidation Effects](#)

[Introduction to Ion Implantation](#)

[Concept of Ion Implantation](#)

[Ion-Implantation System](#)

[Implant Species Sources](#)

[Ionization Chamber](#)

[Mass Analyzing or Ion Selection](#)

[Acceleration Tube](#)

[Wafer Charging](#)

[Beam Focus](#)

[Neutral Beam Trap](#)

[Beam Scanning](#)

[End Station and Target Chamber](#)

[Ion-Implant Masks](#)

[Dopant Concentration in Implanted Regions](#)

[Crystal Damage](#)

[Annealing and Dopant Activation](#)

[Channeling](#)

[Evaluation of Implanted Layers](#)

[Uses of Ion Implantation](#)

[The Future of Doping](#)

[Review Topics](#)

[References](#)

12 Layer Deposition

[Introduction](#)

[Film Parameters](#)

[Chemical Vapor Deposition Basics](#)

[Basic CVD System Components](#)

[CVD Process Steps](#)

[CVD System Types](#)

[Atmospheric-Pressure CVD Systems](#)

[Horizontal-Tube Induction-Heated APCVD](#)

[Barrel Radiant-Induction-Heated APCVD](#)

[Pancake Induction-Heated APCVD](#)

[Continuous Conduction-Heated APCVD](#)

[Horizontal Conduction-Heated APCVD](#)

[Low-Pressure Chemical Vapor Deposition](#)

[Horizontal Conduction-Convection-Heated LPCVD](#)

[Ultra-High Vacuum CVD](#)

[Plasma-Enhanced CVD \(PECVD\)](#)

[High-Density Plasma CVD](#)

[Atomic Layer Deposition](#)

[Vapor-Phase Epitaxy](#)

[Molecular Beam Epitaxy](#)

[Metalorganic CVD](#)

[Deposited Films](#)

[Deposited Semiconductors](#)

[Epitaxial Silicon](#)

[Polysilicon and Amorphous Silicon Deposition](#)

[SOS and SOI](#)

[Gallium Arsenide on Silicon](#)

[Insulators and Dielectrics](#)

[Silicon Dioxide](#)

[Doped Silicon Dioxide](#)

[Silicon Nitride](#)

[High-k and Low-k Dielectrics](#)

[Conductors](#)

[Review Topics](#)

[References](#)

[13 Metallization](#)

[Introduction](#)

[Deposition Methods](#)

[Single-Layer Metal Systems](#)

[Multilevel Metal Schemes](#)

Conductors Materials

Aluminum

Aluminum-Silicon Alloys

Aluminum-Copper Alloy

Barrier Metals

Refractory Metals and Refractory Metal Silicides

Plugs

Sputter Deposition

Copper Dual-Damascene Process

Low-k Dielectric Materials

The Dual-Damascene Copper Process

Barrier or Liner Deposition

Seed Deposition

Electrochemical Plating

Chemical-Mechanical Processing

CVD Metal Deposition

Doped Polysilicon

CVD Refractory Deposition

Metal-Film Uses

MOS Gate and Capacitor Electrodes

Backside Metallization

Vacuum Systems

Dry Mechanical Pumps

Turbomolecular Hi-Vac Pumps

Review Topics

References

14 Process and Device Evaluation

Introduction

Wafer Electrical Measurements

Resistance and Resistivity

Resistivity Measurements

Four-Point Probe

Process and Device Evaluation

Sheet Resistance

Four-Point Probe Thickness Measurement

Concentration or Depth Profile

Secondary Ion Mass Spectrometry

Physical Measurement Methods

Layer Thickness Measurements

Color

- [Spectrophotometers or Reflectometry](#)
- [Ellipsometers](#)
- [Stylus \(Surface Profilometers\)](#)
- [Photoacoustic](#)
- [Gate Oxide Integrity Electrical Measurement](#)
- [Junction Depth](#)
 - [Groove and Stain](#)
 - [Scanning Electron Microscope Thickness Measurement](#)
 - [Spreading Resistance Probe](#)
 - [Secondary Ion Mass Spectrometry](#)
 - [Scanning Capacitance Microscopy](#)
 - [Critical Dimensions and Line-Width Measurements](#)
 - [Optical Image-Shearing Dimension Measurement](#)
 - [Shape Metrology and Optical Critical Dimension](#)
- [Contamination and Defect Detection](#)
 - [1× Visual Surface Inspection Techniques](#)
 - [1× Collimated Light](#)
 - [1× Ultraviolet](#)
 - [Microscope Techniques](#)
 - [Automated In-Line Defect Inspection Systems](#)
- [General Surface Characterization](#)
 - [Atomic Force Microscopy](#)
 - [Scattrometry](#)
- [Contamination Identification](#)
 - [Auger Electron Spectroscopy](#)
 - [Electron Spectroscopy for Chemical Analysis](#)
 - [Time of Flight Secondary Ion Mass Spectrometry](#)
 - [Evaluation of Stack Thickness and Composition](#)
- [Device Electrical Measurements](#)
 - [Equipment](#)
 - [Resistors](#)
 - [Diodes](#)
 - [Bipolar Transistors](#)
 - [MOS Transistors](#)
 - [Capacitance-Voltage Profiling](#)
 - [Device Failure Analysis—Emission Microscopy](#)
- [Review Topics](#)
- [References](#)

15 The Business of Wafer Fabrication

- [Introduction](#)

- [Moore's Law and the New Wafer-Fabrication Business](#)
- [Wafer-Fabrication Costs](#)
 - [Overhead](#)
 - [Materials](#)
 - [Equipment](#)
 - [Labor](#)
 - [Production Cost Factors](#)
 - [Yield](#)
 - [Yield Improvements](#)
 - [Yield and Productivity](#)
 - [Book-to-Bill Ratio](#)
 - [Cost of Ownership](#)
- [Automation](#)
 - [Process Automation](#)
- [Wafer-Loading Automation](#)
- [Clustering](#)
- [Wafer-Delivery Automation](#)
 - [Closed-Loop Control-System Automation](#)
- [Factory-Level Automation](#)
- [Equipment Standards](#)
 - [Fab Floor Layout](#)
 - [Batch versus Single-Wafer Processing](#)
 - [Green Fabs](#)
- [Statistical Process Control](#)
- [Inventory Control](#)
 - [Just-in-Time Inventory Control](#)
- [Quality Control and Certification—ISO 9000](#)
- [Line Organization](#)
- [Review Topics](#)
- [References](#)

16 Introduction to Devices and Integrated Circuit Formation

- [Introduction](#)
- [Semiconductor-Device Formation](#)
 - [Resistors](#)
 - [Capacitors](#)
 - [Diodes](#)
 - [Transistors](#)
 - [Field-Effect Transistors](#)
- [Alternatives to MOSFET Scaling Challenges](#)
 - [Conductors](#)

[Integrated-Circuit Formation](#)

[Bipolar Circuit Formation](#)

[MOS Integrated Circuit Formation](#)

[Bi-MOS](#)

[Silicon on Insulator Isolation](#)

[Superconductors](#)

[Microelectromechanical Systems](#)

[Strain Gauges](#)

[Batteries](#)

[Light-Emitting Diodes](#)

[Optoelectronics](#)

[Solar Cells](#)

[Temperature Sensing](#)

[Acoustic Wave Devices](#)

[Review Topics](#)

[References](#)

17 Process and Device Evaluation

[Introduction](#)

[Circuit Basics](#)

[Integrated Circuit Types](#)

[Logic Circuits](#)

[Memory Circuits](#)

[Redundancy](#)

[The Next Generation](#)

[Review Topics](#)

[References](#)

18 Packaging

[Introduction](#)

[Chip Characteristics](#)

[Package Functions and Design](#)

[Substantial Lead System](#)

[Physical Protection](#)

[Environmental Protection](#)

[Heat Dissipation](#)

[Common Package Parts](#)

[Cleanliness and Static Control](#)

[Basic Bonding Processes](#)

[Wire Bonding Process](#)

[Prebonding Wafer Preparation](#)

[Die Separation](#)

[Die Pick and Place](#)

[Die Inspection](#)

[Die Attach](#)

[Wire Bonding](#)

[Tape Automated Bonding Process](#)

[Bump or Ball Flip-Chip Bonding](#)

[Example Bump or Ball Process](#)

[Copper Metallization \(Damascene\) Bump Bonding
Reflow](#)

[Die Separation and Die Pick and Place](#)

[Alignment of Die to Package](#)

[Attachment to Package \(or Substrate\)](#)

[Deflux](#)

[Underfillment](#)

[Encapsulation](#)

[Postbonding and Preseal Inspection](#)

[Sealing Techniques](#)

[Lead Plating](#)

[Plating Process Flows](#)

[Lead Trimming](#)

[Deflashing](#)

[Package Marking](#)

[Final Testing](#)

[Environmental Tests](#)

[Electrical Testing](#)

[Burn-In Tests](#)

[Package Design](#)

[Metal Cans](#)

[Pin Grid Arrays](#)

[Ball-Grid Arrays or Flip-Chip Ball-Grid Arrays](#)

[Quad Packages](#)

[Thin Packages](#)

[Chip-Scale Packages](#)

[Lead on Chip](#)

[Three-Dimensional Packages](#)

[Stacking Die Techniques](#)

[Three-Dimensional Enabling Technologies](#)

[Hybrid Circuits](#)

[Multichip Modules](#)

[The Known Good Die Problem](#)

[Package Type or Technology Summary](#)

[Package or PCB Connections](#)

[Bare Die Techniques and Blob Top](#)

[Review Topics](#)

[References](#)

[**Glossary**](#)

[**Index**](#)

Preface

From the Preface of the *First Edition*: “As the semiconductor industry becomes more important in the economy, more people will be involved in the industry. It is my intention that *Microchip Fabrication* will serve their needs.”

Indeed the semiconductor industry has grown into a major international industrial segment. The semiconductor materials and equipment industries have also grown into major industrial sectors. This edition has followed the goal of the *First Edition* to serve the training needs of wafer-fabrication workers, whether they be production workers, technicians, professionals in the materials and equipment sectors, or engineers.

The *Sixth Edition* retains the physics, chemistry, and electronic fundamentals underlying the sophisticated manufacturing materials and processes of the modern semiconductor industry. It goes on to profile the state-of-the-art processes that have grown from the simple laboratory productions lines of the 1960s. Not every individual process flow can be detailed in an introductory text. But current technologies used in the patterning, doping, and layering steps are explained. The intention of this book is that the reader will gain enough general knowledge to be able to keep abreast of new processes and equipment.

I am indebted to the valuable input from Anne Miller and Dr. Michael Hynes at Semiconductor Services, Bill Moffat the founder and President of Yield Engineering Systems, and Don Keenan, process engineer extraordinaire.

Kudos to Senior Editor Michael McCabe and his staff at McGraw-Hill for their support and guidance. And a thanks to Sheena Uprety, Associate Project Manager at Cenveo Publisher Services, and the copyeditor, Ragini Pandey, for turning my manuscript into a ready-for-production text.

And, of course, a shout out to my ever supportive and patient wife, Mary DeWitt. She edited the first edition, has given me encouragement during the writing of every edition, and has lent her eagle eye to this latest edition.

Note to Instructors: If you are an instructor using this book as a textbook, then there is an Instructor’s Manual available at www.mhprofessional.com/mf6e.

Peter Van Zant

CHAPTER 1

The Semiconductor Industry

Introduction

In this chapter, you will be introduced to the semiconductor industry with a description of the historic product and process developments that gave rise to a major world industry. The major manufacturing stages, from material preparation to packaged product, will also be introduced along with the mainstream product types, transistor building structures, and the different integration levels. Industry product and processing trends will be identified.

As the industry moved from small-scale laboratory production to vast automated factories, the industry drivers and economics changed. Large specialty materials and equipment industries have developed to support chip manufacturing. Global semiconductors are a \$300 billion industry, and it feeds a \$1.2 trillion global-electronics–systems industry. Going forward, nanotechnology and the explosion of the worldwide consumer markets are shaping the future of the semiconductor industry in ways that are still unfolding.¹ Wafer fabrication has spawned an equipment industry of some \$60 billion in annual sales (typically 15–20% of chip sales).

Birth of an Industry

The electronics industry got its jump start with the discovery of the audion vacuum tube in 1906 by Lee DeForest.² This discovery made the existence of radio, television, and other consumer electronics possible. It also was the brains of the world's first electronic computer, named the Electronic Numeric Integrator and Calculator (ENIAC), which was first demonstrated at the Moore School of Engineering in Pennsylvania in 1947.

The ENIAC was unlike a modern computer. It occupied some 1500 ft², weighed 30 tons, generated large quantities of heat, needed a small power station, and cost \$400,000 in 1940. The ENIAC was based on 19,000 vacuum tubes along with thousands of resistors and capacitors ([Fig. 1.1](#)).

Size, ft	30 × 50
Weight, tons	30
Vacuum Tubes	18,000
Resistors	70,000
Capacitors	10,000
Switches	6000
Power Requirements, W	150,000
Cost (in 1940)	\$400,000

FIGURE 1.1 ENIAC statistics. (Source: Foundations of Computer Technology, J. G. Giarratano, Howard W. Sams & Co., Indianapolis, IN, 1983.)

A vacuum tube consists of three elements: two electrodes separated by a grid in a glass enclosure ([Fig. 1.2](#)). Inside the enclosure is a vacuum, required to prevent the elements from burning up and to allow the easy transfer of electrons.

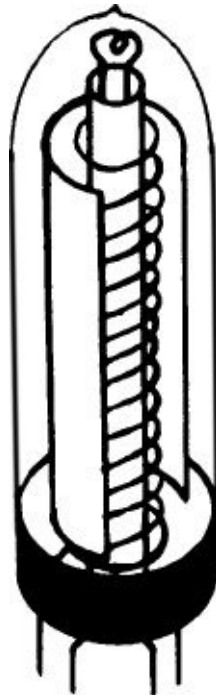


FIGURE 1.2 Vacuum tube.

Tubes perform two important electrical functions: *switching* and *amplification*. Switching refers to the ability of an electrical device to turn a current on or off. Amplification is a little more complicated. It is the ability of a device to receive a small signal (or current) and amplify it while retaining its electrical characteristics.

Vacuum tubes suffer from a number of drawbacks. They are bulky and prone to loose connections and vacuum leaks; they are fragile; they require relatively large amounts of power to operate; and their elements deteriorate rather rapidly. One of the major drawbacks of the ENIAC and other tube-based computers was a limited

operating time due to tube burn out. However, the world did not recognize the potential of computers early on. IBM Chairman, Thomas Watson, in 1943, ventured that, “I think there is a worldwide market for maybe five computers.”

These problems were the impetus leading many laboratories around the country to seek a replacement for the vacuum tube. That effort came to fruition on December 23, 1947, when three Bell Lab scientists demonstrated an electrical amplifier formed from the semiconducting material germanium ([Fig. 1.3](#)).



FIGURE 1.3 The first transistor.

This device offered the electrical functioning of a vacuum tube, but added the advantages of being solid state (no vacuum), being small and lightweight, and having low power requirements and a long lifetime. First named a *transfer resistor*, the new device soon became known as the *transistor*.

The three scientists, John Bardeen, Walter Brattin, and William Shockley were awarded the 1956 Nobel Prize in physics for their invention.

The Solid-State Era

That first transistor was a far distance from the high-density integrated circuits of today. But it was the component that gave birth to the solid-state electronics era with all its famous progeny. Besides transistors, solid-state technology is used to create diodes, resistors, and capacitors. Diodes are two-element devices that function in a circuit as an on/off switch. Resistors are mono-element devices that serve to limit

current flow. Capacitors are two-element devices that store charge in a circuit. In some integrated circuits, the technology is used to create fuses. Refer to [Chap. 14](#) for an explanation of these concepts and an explanation of how these devices work.

These devices, containing only one device per chip, are called *discrete* devices ([Fig. 1.4](#)). Most discrete devices have less-demanding operational and fabrication requirements than integrated circuits. In general, discrete devices are not considered leading-edge products. Yet, they are required in most sophisticated electronic systems. In 1998, they accounted for 12 percent of the dollar volume of all semiconductor devices sold.³ The semiconductor industry was in full swing by the early 1950s, supplying devices for transistor radios and transistor-based computers.

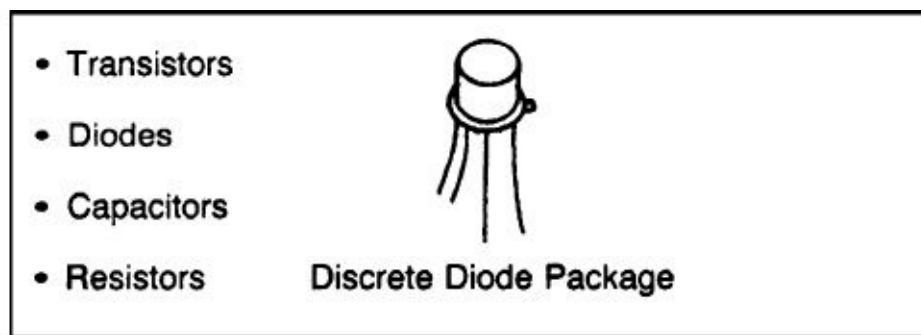


FIGURE 1.4 Solid-state discrete devices.

Integrated Circuits (ICs)

The dominance of discrete devices in solid-state circuits came to an end in 1959. In that year, Jack Kilby, a new engineer at Texas Instruments in Dallas, Texas, formed a complete circuit on a single piece of the semiconducting material germanium. His invention combined several transistors, diodes, and capacitors (five components total) and used the natural resistance of the germanium chip (called a *bar* by Texas Instruments) as a circuit resistor. This invention was the *integrated circuit*, the first successful integration of a complete circuit in and on the same piece of a semiconducting substrate.

The Kilby circuit did not have the form that is prevalent today. It took Robert Noyce, then at Fairchild Camera, to furnish the final piece of the puzzle. In [Fig. 1.5](#) is a drawing of the Kilby circuit. Note that the devices are connected with individual wires.

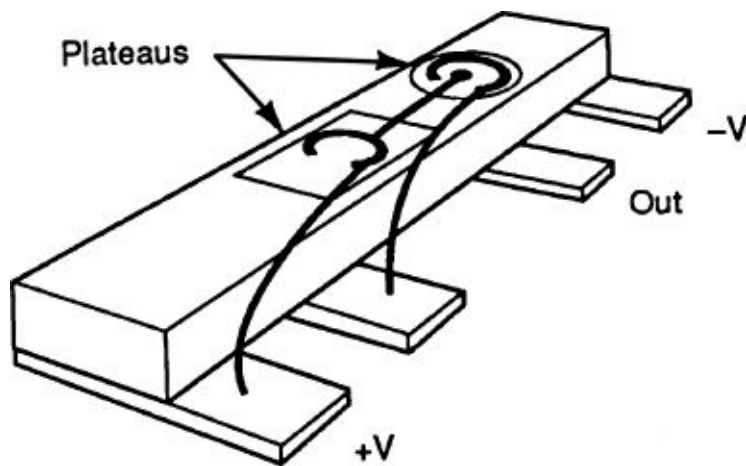


FIGURE 1.5 Kilby integrated circuit from his notebook.

Earlier, Jean Horni, also at Fairchild Camera, had developed a process of forming electrical junctions in the surface of a chip to create a solid-state transistor with a flat profile ([Fig. 1.6](#)). The flattened profile was the outcome of taking advantage of the easily formed natural oxide of silicon, which also happened to be a dielectric (electrical insulator). Horni's transistor used a layer of evaporated aluminum, which was patterned into the proper shape, to serve as wiring for the device. This technique is called *planar technology*. Noyce applied this technique to *wire* together the individual devices previously formed in the silicon wafer surface ([Fig. 1.7](#)).

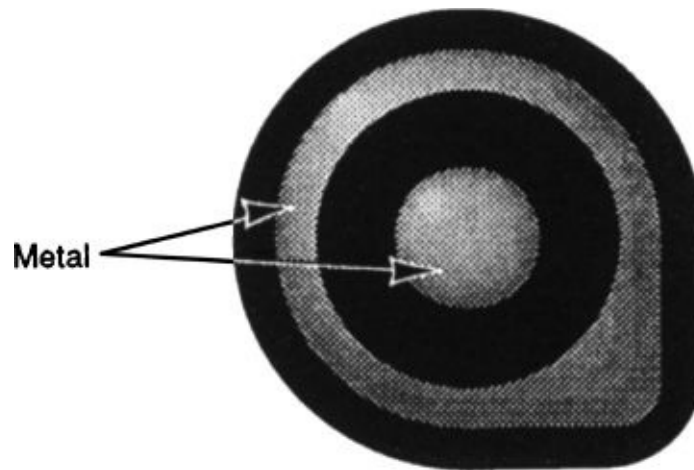


FIGURE 1.6 Horni *teardrop* transistor.

NOYCE PATENT U.S. PATENT No. 2,981,877

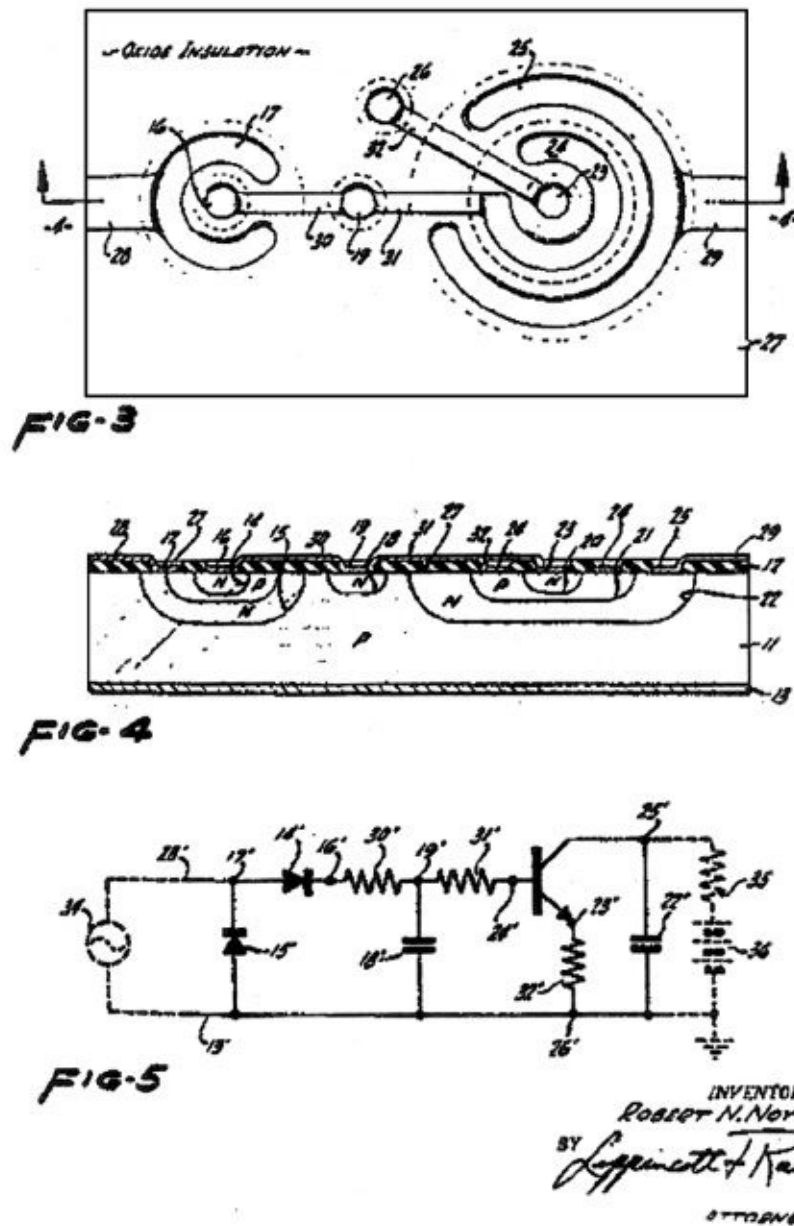


FIGURE 1.7 Noyce IC patent. (Courtesy of Semiconductor Reliability News, June 2003.)

The Noyce integrated circuit became *the* model for all integrated circuits. The techniques used not only met the needs of that era, but contained the seeds for all the miniaturization and cost-effective manufacturing that still drives the industry. Kilby and Noyce shared the patent for the integrated circuit.

Process and Product Trends

Since 1947, the semiconductor industry has seen the continuous development of new and improved processes. These process improvements have in turn led to the more

highly integrated and reliable circuits that have, in their turn, fueled the continuing electronics revolution. These process improvements fall into two broad categories: process and structure. *Process improvements* are those that allow the fabrication of the devices and circuits in smaller dimensions, in ever higher density, quantity, and reliability. The structure improvements are the invention of new device designs allowing greater circuit performance, power control, and reliability.

Device component size and the number of components in an IC are the two common trackers of IC development. Component dimensions are characterized by the smallest dimension in the design. This is called the *feature size* and is usually expressed in microns or nanometers. A micron is 1/1,000,000 of a meter or about 1/100 the diameter of a human hair. A nanometer is 1/1,000,000,000 of a meter. A more specific tracker of semiconductor devices is *gate width*. Transistors are composed of three parts, one of which acts to allow the passage of current. In today's technology, the most popular transistor is the metal-oxide-semiconductor field effect transistor (MOSFET) structure ([Chap. 16](#)). The controlling part is called the *gate*. Smaller gate widths drive the industry by producing smaller and faster transistors and more dense circuits. Currently, the industry is driving to the 5-nm gate width, with projections in the *International Technology Roadmap for Semiconductors* projecting a 5-nm gate width around 2016.⁴

Moore's Law

In 1965, Gordon Moore, a founder of Intel, noted that the number of transistors on a chip were doubling annually. He published the observation, which was immediately dubbed *Moore's law*. He updated the law to a doubling every two years. Industry observers have used this law to predict the future density of chips. Over the years, it has proven very accurate and now drives technical advances. If it holds true, the transistor count on a chip could reach into the billions ([Fig. 1.8](#)). It is the basis of the *International Technology Roadmap for Semiconductors (ITRS)*, developed by the Semiconductor Industry Association (SIA).

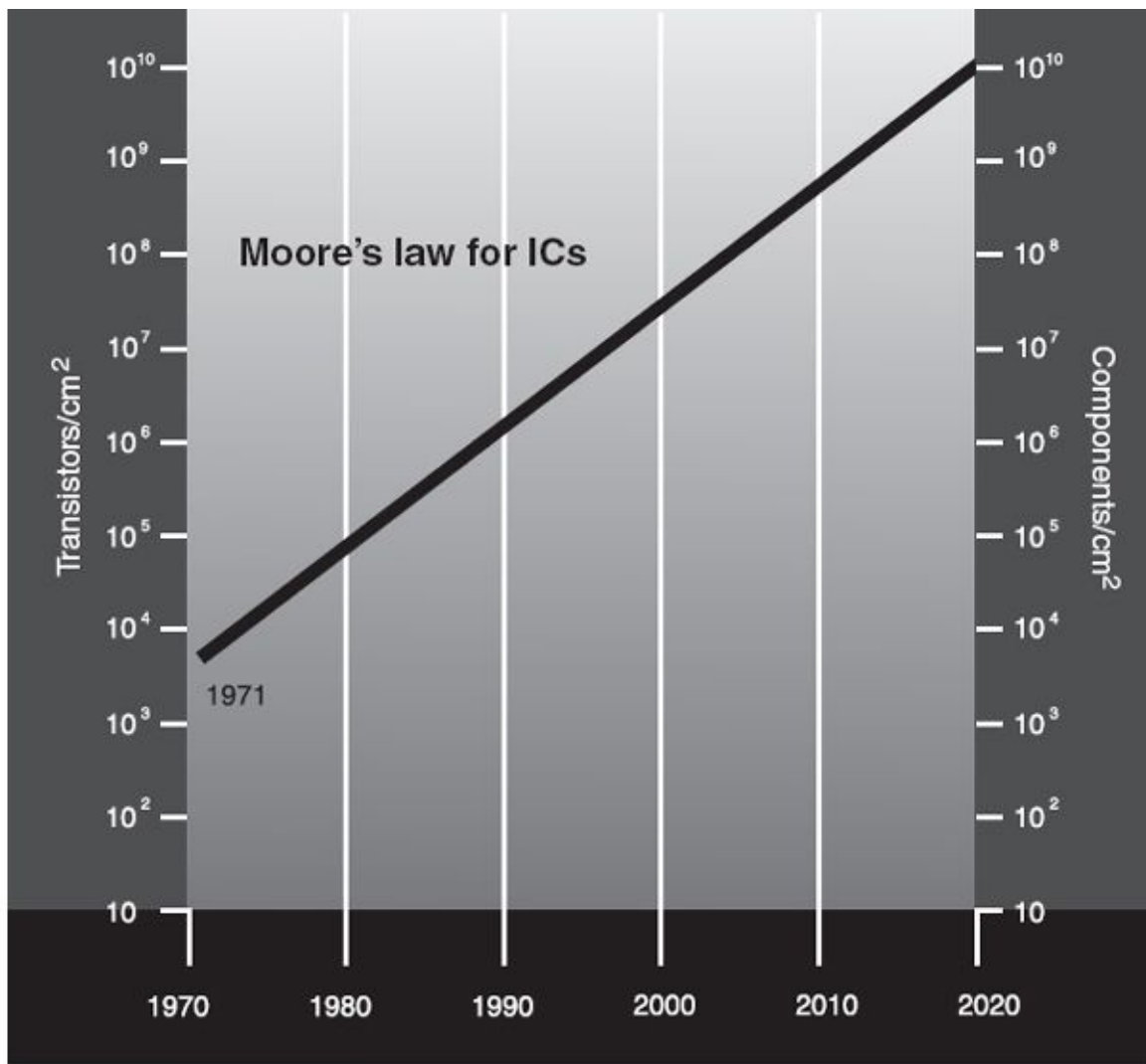


FIGURE 1.8 Moore's law. (Source: *Moore's Law Meets Its Match*, IEEE Spectrum, June 2006.)

There is speculation that chip density may exceed the Moore's law projection. Georgia Tech's Microsystems Packaging Research notes that from about 50 components per square centimeter in 2004, component density will climb to about a million per square centimeter by 2020.⁵

The continued increase in the component density in a chip has indeed followed Moore's law. There is also a discussion that the industry has now adapted Moore's law as the future driver (goals) of chip density and improved performance. These goals are embedded in the latest editions of the *International Technology Roadmap for Semiconductors*, 2011 Edition.

Circuit density is tracked by the *integration level*, which is the number of components in a circuit. Integration levels (Fig. 1.9) range from small-scale integration (SSI) to ultra large-scale integration (ULSI). ULSI chips are sometimes referred to as very very large-scale integration (VVLSI). The popular press calls these newest products *megachips*.

Level	Abbreviation	# Components per Chip
Small-Scale Integration	SSI	2 - 50
Medium Scale Integration	MSI	50 - 5000
Large-Scale Integration	LSI	5000 - 100,000
Very Large Scale Integration	VLSI	Over 100,000 - 1,000,000
Ultra Large Scale Integration	ULSI	> 1,000,000

FIGURE 1.9 IC integration table.

In addition to the integration scale, memory circuits are identified by the number of memory bits contained in the circuit (a four-meg memory chip can store four million bits of memory). Logic circuits are often rated by the number of *gates* they have. A gate is the basic operational component of a logic circuit.

Decreasing Feature Size

The journey from small-scale integration to today's megachips has been driven primarily by reductions in the feature size of the individual components. This decrease has been brought about by dramatic increases in the imaging process, known as lithography, and the trend to multiple layers of conductors. The Semiconductor Industry Association has projected feature sizes decreasing to 22 nm (0.0022 μm) by the year 2016.⁶ Along with the ability to make components on the chip smaller comes the benefit of crowding them closer together, further increasing density ([Fig. 1.10](#)).

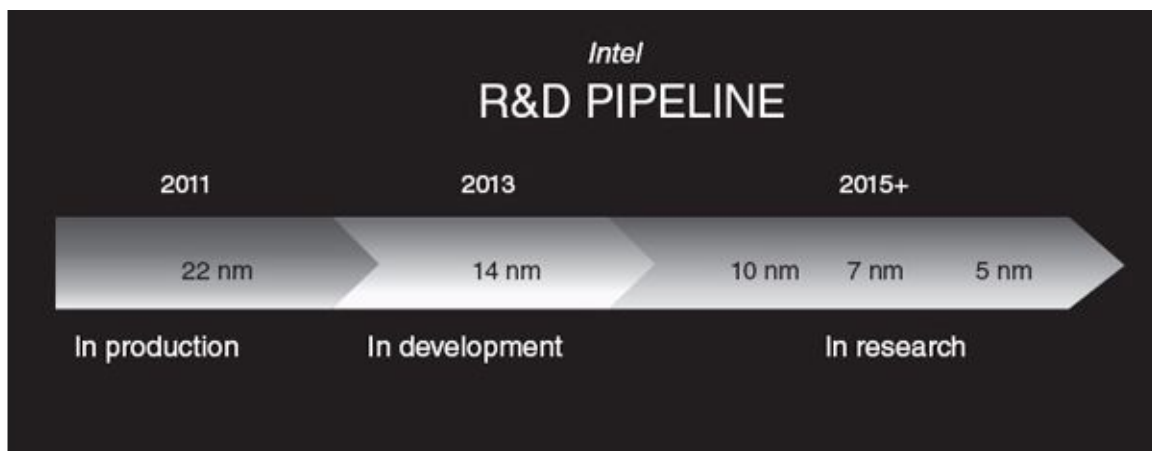


FIGURE 1.10 Intel feature size projection.

An analogy used to explain these trends is the layout of a neighborhood of single-family homes. The density of the neighborhood is a function of the house size, lot size, and the width of the streets. Accommodating a higher population could come by increasing the size of the neighborhood (increasing the chip area). Another

possibility is to reduce the size of the individual houses and place them on smaller lots. We can also reduce the street size to increase density. However, at some point, the streets cannot be reduced anymore in size or they would not be wide enough for autos. Furthermore, at some point, the houses cannot be further reduced in size and still function as dwelling units. At this point, an option is to build up by building multidecked freeways and/or replacing individual homes with apartment buildings. All of these concepts are used in semiconductor technology.

There are several benefits to the reduction of the feature size and its attendant increase in circuit density. At the circuit performance level, there is an increase in circuit speed. With lesser distances to travel and with the individual devices occupying less space, information can be put into and gotten out of the chip in lesser time. Anyone who has waited for their personal computer to perform a simple operation can appreciate the effect of faster performance. These same density improvements result in a chip or circuit that requires less power to operate. The small power station required to run the ENIAC has given way to powerful laptop computers that run on a set of batteries.

Increasing Chip and Wafer Size

The advancement of chip density from the SSI level to ULSI chips has driven larger chip sizes. Discrete and SSI chips average about 100 mil (0.1 in) on a side. ULSI chips are in the 500 to 1000 mil (0.5 to 1.0 in) per side, or larger, range. ICs are manufactured on thin disks of silicon (or other semiconductor material, see [Chap. 2](#)) called *wafers*. Placing square or rectangular chips on a round wafer leave unavailable areas around the edge (see [Fig. 6.6](#)). These unavailable areas can become large as the chip size increases ([Fig. 1.11](#)). The desire to offset the loss of usable silicon has driven the industry to larger wafers. As the chip size increases, the 1-in diameter wafers of the 1960s have given way to 200- and 300-mm (8- and 12-in) sized wafers. Production efficiency increases, because the area of a circle increases as the mathematical square of the radius. Thus, doubling the wafer diameter from 6 to 12 in increases the area available for chip fabrication by four times.

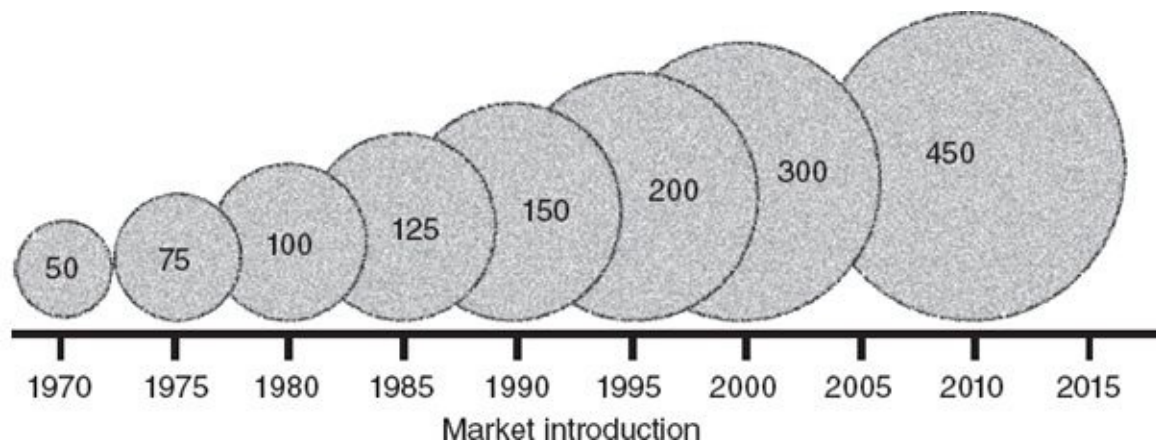


FIGURE 1.11 Wafer size history. (Courtesy of Future Fab International.)

The projected year for the introduction of 450-mm (18-in) diameter wafers was 2012. Despite another recession, 450-mm wafers became available and Intel, TSMC, and Samsung announced plans to build new wafer fabrication production plants (fabs). Cost has been a primary barrier to the processing of larger wafers. Generally it is not technically possible to simply expand a 300-mm production line. Therefore new fab facilities become necessary, but not before the equipment suppliers design, test, and build expanded capacity process tools. These moves are expensive and time-consuming. But the real outcomes of more efficient production, yields, and accommodation of advanced circuits have driven the industries, continued advances.² The expense factor has also led to the retention (tail) of smaller diameter wafer production lines. For established older product lines being fabricated in plants that have long been paid for, there is little economic incentive to move to larger wafers. Indeed, 150-mm (5.9-in aka 6-inch) wafers, as well as 200-mm wafers are still in use.

Reduction in Defect Density

As feature sizes have decreased, the need for reduced defect density and defect size on the chips (and in the manufacturing process) has become critical. A 1- μm piece of dirt on a 100- μm sized transistor may not be a problem. On a 1- μm sized transistor, it becomes a killer defect that can render the component inoperable ([Fig. 1.12](#)). Contamination control programs have become a must have for successful microchip fabrication (see [Chap. 5](#)).

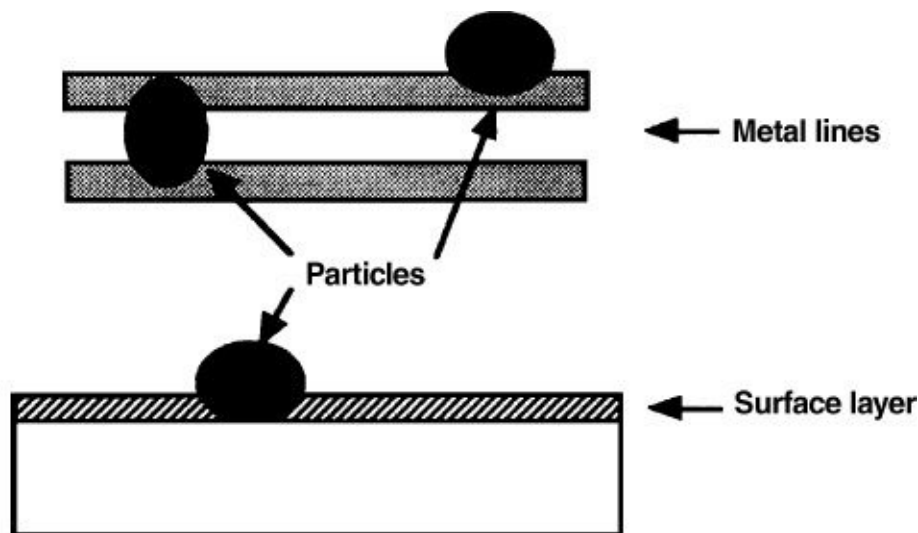


FIGURE 1.12 Relative size of airborne particles and wafer dimensions.

Increase in Interconnection Levels

The component density increase has led to a *wiring* problem. In the neighborhood analogy, reducing street widths was one strategy to increase density. But, at some point, the streets become too narrow to allow cars to travel. The same thing happens in IC design. The increased component density and close packing rob the surface space needed on the surface to connect the components. The solution is multiple levels of *wiring* stacked ([Fig. 1.13](#)) above the surface components in layers of insulators and conducting layers ([Chap. 13](#)).

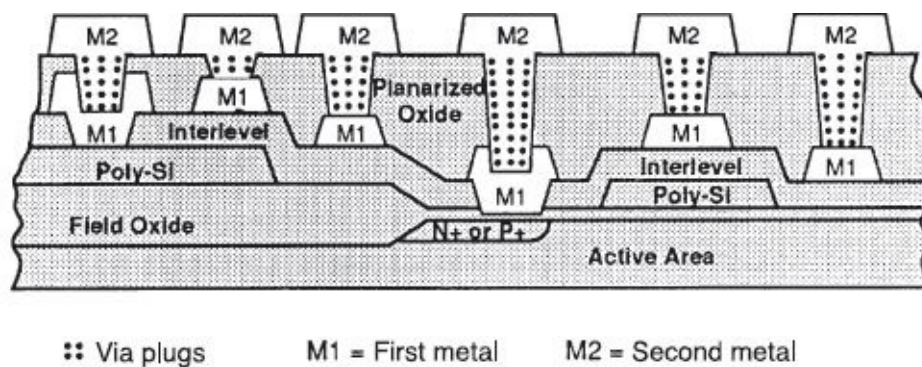


FIGURE 1.13 Cross-section of typical planarized two-level metal VLI structure showing range of via depths after planarization. (*Courtesy of Solid State Technology.*)

Planarization is the formation of the active transistors and other components in the base wafer (usually silicon). In 2011, Intel Corporation announced a new three-dimensional (3D) device ([Fig. 1.14](#)) with the active transistor gate sticking above the wafer.⁸ The device was called a tri-gate transistor. The device performance is enhanced from the increased gate surface area ([Chap. 16](#)).

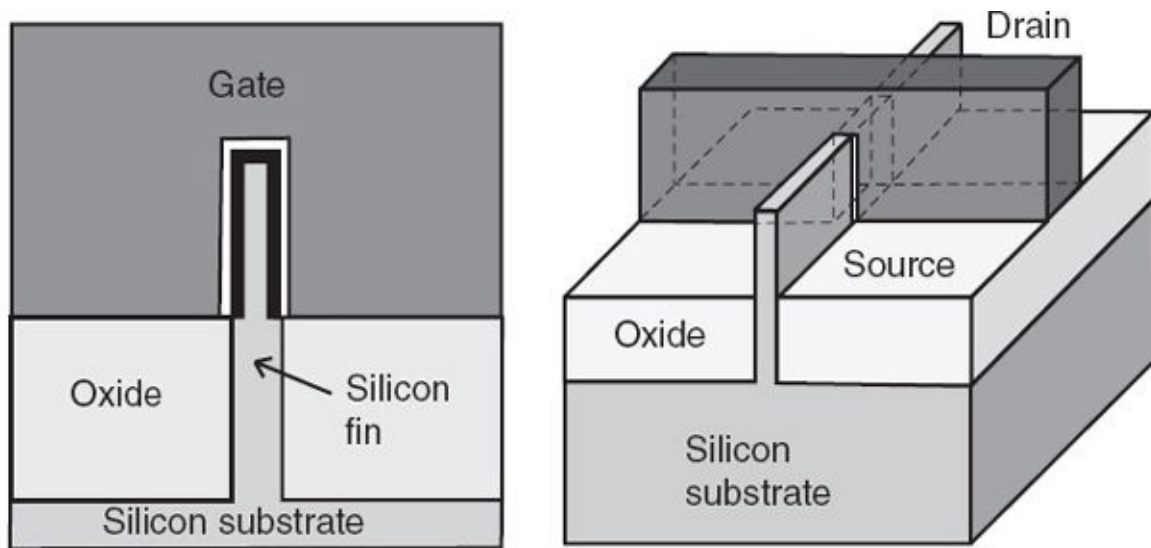


FIGURE 1.14 Tri-gate transistor goes 3D as Intel reinvents the microchip.

The Semiconductor Industry Association Roadmap

These major IC parameters are interrelated. Moore's law predicts the future of component density, which triggers the calculation of the integration level (component density), chip size, defect density (and size), and the number of interconnection levels required. The SIA and partners have made these projections into the future through the *International Technology Roadmap for Semiconductors* in a series of roadmaps covering these and other critical device and production parameters. In addition to projecting component, process, and wafer parameters, it identifies future performance standards for companion materials and equipment needed to support the advanced components.

Chip Cost

Perhaps the most significant effect of these process and product improvements is the cost of the chips. The reductions are typical for any maturing product. Prices start high, and as the technology is mastered and manufacturing efficiencies increase the prices drop and eventually become stable. These chip prices have constantly declined even as the performance of the chips has increased. The factors affecting chip cost are discussed in [Chap. 15](#).

The two factors, increased performance and less cost, have driven the explosion of products using solid-state electronics. By the 1990s, an auto had more computing power onboard than the first lunar space shots. Even more impressive is the personal computer. Today, for a moderate price, a desktop computer can deliver more power than an IBM mainframe manufactured in 1970. Major industry use of chips is shown in [Fig. 1.15](#).

Year/System	Consumer	Auto	Computer	Industrial	Commercial	Gov/Mil
2010	29%	1.70%	34.20%	4.90%	29.90%	0.20%
2016 (Projected)	14.40%	3.20%	43.10%	3.70%	35.40%	0.20%

FIGURE 1.15 Flash memory usage by system. (Adapted from IC Insights—Market Drivers 2013 Products and Services, Scotsdale, AZ.)

The history of the semiconductor industry is one of the continual developments and advances emerging to world dominance in the mid-1990s. In that decade, the semiconductor industry became the nation's leading value-added industry, outperforming the auto industry.

Industry Organization

The electronics industry is divided into two major segments: semiconductors and systems (or products). The *Semiconductor* segment encompasses the material suppliers, circuit design, chip manufacturers, and all of the equipment and chemical suppliers to the industry. The systems segment encompasses the industry that designs and produces the vast number of semiconductor-device-based products, from consumer electronics to space shuttles. The electronics industry includes the manufacturers of printed circuit boards.

The semiconductor segment is composed of two major subsegments. One includes the firms that actually make the semiconductor solid-state devices and circuits. Within this segment, there are three types of chip suppliers: integrated device manufacturers (IDMs) design, manufacture, package, and market chips. Foundry companies build circuit chips for other chip suppliers. Waferless (or fabless) companies design and market chips, buying finished chips from chip foundries. Chips are fabricated by both merchant and captive producers. Merchant suppliers manufacture just chips and sell them on the open market. Captive suppliers are firms whose final product is a computer, communications system, or other product, and they produce chips in house for their own products. Some firms produce chips for in-house use and also sell on the open market, and others produce specialty chips in house and buy others on the open market. Since the 1980s, the trend has been to a greater percentage of chips being fabricated in captive fab areas.

Stages of Manufacturing

Solid-state devices are manufactured in the following five distinct stages ([Fig. 1.16](#)):

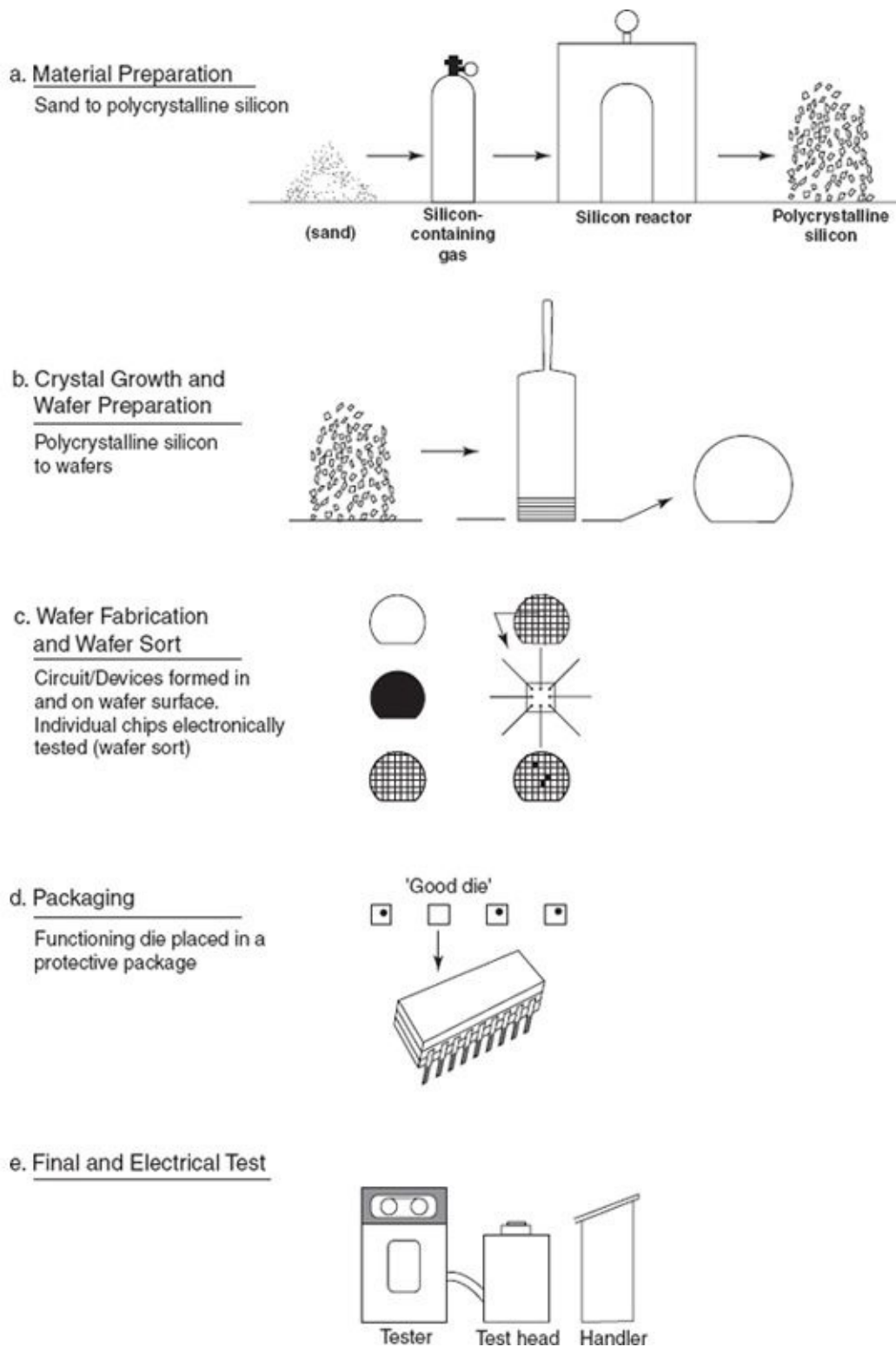


FIGURE 1.16 Stages of semiconductor production.

1. Material preparation
2. Crystal growth and wafer preparation
3. Wafer fabrication and sort

4. Packaging

5. Final and electrical tests

In the first stage, material preparation (see [Chap. 2](#)), the raw semiconducting materials are mined and purified to meet semiconductor standards. For silicon, the starting material is sand, which is converted to pure silicon with a polysilicon structure ([Fig. 1.16a](#)).

In stage two, the material is formed into a crystal with specific electrical and structural parameters. Next, thin disks called *wafers* are cut from the crystal and surface-treated ([Fig. 1.16b](#)) in a process called crystal growth and wafer preparation (see [Chap. 3](#)). The industry also makes devices and circuits from germanium and compounds of different semiconductor materials.

In stage three ([Fig. 1.16c](#)), wafer fabrication, the devices or integrated circuits are actually formed in and on the wafer surface. Up to several thousand identical devices can be formed on each wafer, although 200 to 300 is a more common number. The area on the wafer occupied by each discrete device or integrated circuit is called a *chip* or *die*. The wafer fabrication process is also called fabrication, fab, chip fabrication, or microchip fabrication. While a wafer fabrication operation may take several thousand individual steps, there are two major activities. In the front end of the line (FEOL), the transistors and other devices are formed in the wafer surface. In the back end of the line (BEOL), the devices are wired together with metallization processes, and the circuit is protected with a final sealing layer.

Following wafer fabrication, the devices or circuits on the wafer are complete, but untested and still in wafer form. Next comes an electrical test (called *wafer sort*) of every chip to identify those that meet customer specifications. Wafer sort may be the last step in the wafer fabrication or the first step in the *packaging* process.

Packaging ([Fig. 1.16d](#)) is the series of processes that separate the wafer into individual die and place them into protective packages. This stage also includes final testing of the chip for conformance to customer specifications. The industry also refers to this stage as *assembly and test* (A/T). A protective chip package is necessary to protect the chip from contamination and abuse, and to provide a durable and substantial electrical lead system to allow connection of the chip onto a printed circuit board or directly into an electronic product. Packaging takes place in a different department of the semiconductor producer and quite often in a foreign plant.

The vast majority of chips are packaged in individual packages. But a growing percentage are being incorporated into hybrid circuits, in multichip modules (MCMs), 3-D stacks, or mounted directly on printed circuit boards (chip-onboard, COB). An integrated circuit is an electrical circuit formed entirely by semiconductor technology on a single chip. A hybrid circuit combines semiconductor devices (discretes and ICs) with thick-or thin-film resistors and

conductors and other electrical components on a ceramic substrate. These techniques are explained in [Chap. 18](#).

Six Decades of Advances in Microchip Fabrication Processes

While the tremendous advantages of solid-state electronics was recognized early on, the advancements possible from miniaturization were not realized until two decades later. During the 1950s, engineers set to work and defined many of the basic processes and materials still used today.

William Shockley and Bell Labs get much of the credit for the spread of semiconductor technology. Shockley left Bell Labs in 1955 and formed Shockley Laboratories in Palo Alto, California. While his company did not survive, it established semiconductor manufacturing on the West Coast and provided the beginning of what eventually became known as Silicon Valley. Bell Labs helped the fledgling industry with the decision to license its semiconductor discoveries to a host of companies.

The early semiconductor devices were made with the material germanium. Texas Instruments changed that trend with the introduction of the first silicon transistor in 1954. The issue over which material would dominate was settled in 1956 and 1957 by two more developments from Bell Labs: diffused junctions and oxide masking.

It was the development of oxide masking that ushered in the “silicon age.” Silicon dioxide (SiO_2) grows uniformly on silicon and has a similar index of expansion, which allows high-temperature processing without warping. Silicon dioxide is a dielectric material, which allows it to function on the silicon surface as an insulator. Additionally, SiO_2 is an effective block to the dopants that form the N and P regions in silicon.

The net effect of these advances was planar technology introduced by Fairchild Camera in 1960. With the above-named techniques, it was possible to form (diffusion) and protect (silicon dioxide) junctions during and after the wafer-fabrication process. Also, the development of oxide masking allowed two junctions to be formed through the top surface of the wafer that is, in one *plane*. It was this process that set the stage for the development of thin film wiring.

Bell Labs conceived of forming transistors in a high-purity layer of semiconducting material deposited on top of the wafer (see [Chap. 12](#)). This was called an *epitaxial layer*. This discovery allowed higher speed devices and provided a scheme for the closer packing of components in a bipolar circuit.

The 1950s was indeed the golden age of semiconductor development. During this incredibly short time, most of the basic processes and materials were discovered. The decade opened with essentially laboratory level processing, producing small volumes of crude devices in germanium and ended with the first integrated circuit

and silicon firmly established as the semiconductor of the future.

The 1960s was the decade the industry started growing into a sophisticated industry, driven by new products that demanded new fabrication processes, new materials, and new production equipment. The chip price erosion trend of the industry, well established in the 1950s, also was an industry driver.

Technology spread as engineers changed companies in the industry clusters in Silicon Valley, Route 128 around Boston, and in Texas. By the 1960s, the number of fab areas had grown sufficiently, and processes were approaching a level of commonality that attracted semiconductor specialty suppliers.

On the company front, many of the key players of the 1950s formed new companies. Robert Noyce left Fairchild to found Intel (with Andrew Grove and Gordon Moore), and Charles Sporck also left Fairchild to grow National Semiconductor into a major player. Signetics became the first company dedicated exclusively to the fabrication of ICs. New device designs were the usual driver of start-up companies. However, the ever-present price erosion was a cruel trend that drove both established and new companies out of business.

Price dropping was accelerated by the development of a plastic package for silicon devices in 1963. Also in that year, RCA announced the development of the *insulated field effect transistor* (IFET), which paved the way for the MOS industry. RCA also pioneered the first complementary MOS (CMOS) circuits.

At the start of the 1970s, the industry was manufacturing ICs primarily at the MSI level. The move to profitable, high-yield LSI devices was being somewhat hampered by mask-caused defects and the damage inflicted on the wafers by the contact aligners. The mask and aligner defect problem was solved with the development of the first practical projection aligner by Perkin and Elmer Company.

The decade also saw the improvement of cleanroom construction and operation, the introduction of ion implantation machines, and the use of e-beam machines for high-quality mask generation, and mask steppers began to show up in fab areas for wafer imaging.

Automation of processes started with spin-bake and develop-bake systems. The move from operator control to automatic control of the processes increased both wafer throughput and uniformity.

The focus in the 1980s was automation of all phases of wafer fabrication and packaging and elimination of operators from the fab areas. Automation increases manufacturing efficiency, minimizes processing errors, and keeps the wafer fabrication areas less contaminating. Wafers of 300 mm were introduced in the late 1990s, further driving the need for automated fabs ([Chaps. 4 and 15](#)).

The 1980s started with American and European dominance and ended as a worldwide industry. Through the 1970s and 1980s, the 1- μm feature size barrier loomed as both opportunity and challenge. The opportunity was a new era of

megachips with vastly increased speeds and memory. The challenge was the limitations of conventional lithography, additional layers, more step height variation on the wafer surface, and increasing wafer diameters, to mention a few. The 1- μm barrier was crossed in the early 1990s when 50 percent⁷ of the microchip fabrication lines were working at the micron or submicron level.

The industry matured into more traditional focuses on manufacturing and marketing issues. Early on, the profit strategy was to ride the innovation curve. That meant always being first (or close to first) with the latest and greatest chip that could be sold with enough profit to pay for the R&D and finance new designs. The spread of the technology (competition) and improvements in process control, however, moved the industry to greater emphasis on the production issues. The primarily productivity factors are automation, cost control, process characterization and control, and worker efficiency (see [Chap. 15](#)).

Wafer fab facilities are in the gigabuck level (\$3 billion and going up), and equipment and process development are equally expensive. Manufacturing chips with feature sizes below 0.35 μm will require extensive and expensive development of conventional lithography or x-ray and deep UV (DUV) lithography.

The challenge of the *SIA Roadmap* (IRTS) is that many of the processes required to produce the next generations of chips are unknown or in very primitive states of development. However, the good news is that the industry is moving forward along an evolutionary curve rather than relying on revolutionary breakthroughs. Engineers are wringing every bit of productivity out of the processes before looking for a big technology jump to solve problems. This is another sign of a maturing industry.

A major technological change was copper wiring.⁹ Aluminum wiring ran into limitations in several areas, notably in contact resistance with silicon. Copper has always been a better conductor but was difficult to deposit and pattern. It was also a killer of circuit operation if it got into the silicon. IBM developed usable copper processes ([Chaps. 10](#) and [13](#)), which gained almost instant acceptance for wiring together advanced chips.

The Nano Era

Microtechnology in the popular sense means *small*. In the science world, it refers to one-billionth. Thus, feature sizes and gate widths are expressed in microns (micrometers), as in 0.018 μm . It is becoming more common to use nanometers (1×10^{-9} m), thus making the above gate width 180 nm ([Fig. 1.17](#)).¹⁰

	Meters (m)
Meter (m)	1
Centimeter (cm)	1/100
Millimeter (mm)	1/1,000
Micrometer (micron) (μm)	1/1,000,000
Nanometer (nm)	1/1,000,000,000
Angstrom ($\text{\AA}$)	1/10,000,000,000

FIGURE 1.17 Comparative unit lengths.

The way to the nano future is sketched out in the SIA's ITRS. Gate widths of 10 nm or less are predicted by 2016. At these levels, the operational parts of devices consist of only a few atoms or molecules.

Getting there will not be easy. There is a predictable train of events that happen as devices are scaled to smaller dimensions. Advantages are faster operating transistors and higher-density chips. However, smaller dimensions require cleaner environments, increased process control, sophisticated patterning tools, and more.

Wafer diameters are moving to more than 450 mm, and factory automation will be at the tool-to-tool level with onboard process monitoring. More processes at higher levels of detail will require higher-volume wafer fabrication plants with more sophisticated process automation and factory management. Price tags for these mega-plants are headed to the \$10 billion level.¹⁰ This level of investment will pressure faster R&D activities and quick factory start-ups.

By 2016, the industry and circuits will be far different from what they are now, and the industry will be near the end of the basic physics of silicon transistors. Postsilicon production materials are yet to be identified, but the industry will grow. Not all IC uses have to be state of the art. It is unlikely that toasters, refrigerators, and automobiles will require cutting-edge devices. New base materials are in the R&D labs. Compound semiconductors, such as gallium arsenide (GaAs) are candidates.

Another use of the term "nano" is a new way to build very small structures, called *nanotechnology*. It is based on the discovery of a structure of carbon flat crystals shaped like a hollow tube (*nanotube*). These structures have promise for a number of uses. In semiconductor technology, it appears that these nets of carbon atoms can be doped to act as electronic devices and, eventually, electronic circuits and are of interest to solar device manufacturers.

It is safe to say that the semiconductor industry will continue to be the dominant industry as it continues to push the limits of material and manufacturing technology.

It is also safe to predict that the use of ICs will continue to shape our world in ways yet unknown.

Review Topics

Upon completion of this chapter, you should be able to:

1. Describe the difference between discrete devices and integrated circuits.
2. Define the terms solid-state, planar processing, and N-type and P-type semiconducting materials.
3. List the four major stages of semiconductor processing.
4. Explain the integration scale and at least three of the implications of processing circuits of different levels of integration.
5. List the major process and device trends in semiconductor processing.

References

1. McLean, B., "IC Insights," *ISS Kicks Off with IC Industry Reality Talks*, M. A., Fury, *Solid State Technology*, Jan. 16, 2012.
2. Antebi, E., *The Electronic Epoch*, Van Nostrand Reinhold, New York, 1984:126.
3. "Economic Indicator," *Semiconductor International*, Jan. 1998:176.
4. Shankland, S., "Moore's Law: The Rule that Really Matters," *CNET*, Oct. 12, 2012.
5. Tummala, R. R., "Moore's Law Meets Its Match," *IEE Spectrum*, June 2006.
6. Semiconductor Industry Association, *International Technology Roadmap for Semiconductors*, 2001/2003 update, www.semichips.org.
7. Pedus, M. L., "Industry Agrees on First 450-mm Wafer Standard," *EE Times*, Oct. 22, 2008.
8. Stokes, J., "Tri-Gate Transistor from Transistors Go 3-D as Intel ReInvents the Microchip," *ARS Technical*, May 4, 2011.
9. Singer, P., "Copper Goes Mainstream: Low k to Follow," *Semiconductor International*, Nov. 1997:67.
10. Baliga, J. (ed.), *Semiconductor International*, Jan. 1998:15.
11. Flamm, K., "More for Less: The Economic Impact of Semiconductors," Dec. 1997.
12. Hatano, D., "Making a Difference: Careers in Semiconductors,"

Semiconductor Industry Association, Matec Conference, Aug. 1998.

13. SIA, "Economic Indicator," *Semiconductor International*, Jan. 1998:176–177.

14. Rose Associates, Semiconductor Equipment and Materials International (SEMI) Information Seminar, 1994.

15. Skinner, C., and Gettel, G., *Solid State Technology*, Feb. 1998:48.

CHAPTER 2

Properties of Semiconductor Materials and Chemicals

Introduction

Semiconductor materials possess electrical, chemical, and physical properties that allow the unique functions of semiconductor devices and circuits. In this chapter, these properties are examined along with their basics of atoms, electrical classification of solids, and intrinsic and doped semiconductors.

The fabrication of a semiconductor device requires the addition of various layers that perform specific functions in the devices. These materials have specific properties and must be added to the wafer through carefully selected and controlled physical and chemical processes.

The basic properties of gases, acids, bases, and solvents are discussed and illustrated in this chapter. Specialty chemicals are discussed in [Chap. 5](#) (“Contamination Control”) and the specific process chapters.

A unique property of a semiconducting material that sets it apart from metals and dielectric materials is that specific elements can be added, through doping, to change and control its electrical properties. These properties and the results are also described in this chapter. While silicon is the most used semiconducting material, there are others used for their specific properties. Several are identified and discussed in this chapter.

Atomic Structure

The Bohr Atom

The understanding of semiconductor materials requires a basic knowledge of atomic structure.

Atoms are the building blocks of the physical universe. Everything in the universe (as far as we know) is made from the 96 stable materials and 12 unstable ones known as *elements*. Each element has a different atomic structure. The different structures give rise to the different physical and chemical properties of each element.

The unique properties of gold, for example, are due to its atomic structure. If a piece of gold is divided into smaller and smaller pieces, one eventually arrives at the smallest possible piece that identifies gold from other elements. This is the atom of gold. However, it takes a collection of gold atoms before the unique properties are evident.

Dividing that atom further will yield the three parts that compose individual atoms. They are called the *subatomic particles*. These are *protons*, *neutrons*, and

electrons. Each of these subatomic particles has its own properties. A particular combination and structure of the subatomic particles is required to form the gold atom. The basic structure of the atom most used to understand physical, chemical, and electrical differences is attributed to the famous physicist Niels Bohr ([Fig. 2.1](#)).

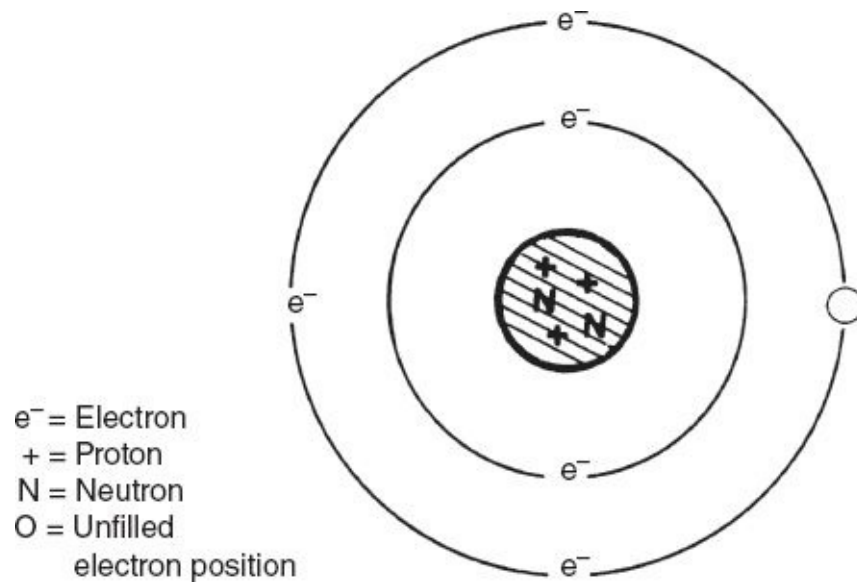


FIGURE 2.1 Bohr atom model.

The Bohr atom model has the positively charged protons and neutral neutrons located together in the nucleus of the atom. The negatively charged electrons move in defined orbits about the nucleus, similar to the movement of the planets about the sun. There is an attractive force between the positively charged protons and the negatively charged electrons. This force is balanced by the outward centrifugal force of the electrons moving in their orbits. The net result is a structurally stable atomic structure.

Each orbit has a maximum number of positions available for electrons. In some atoms, not all of the positions are filled, leaving a hole in the structure. When a particular electron orbit is filled to the maximum, additional electrons must go into the next outer orbit.

The Periodic Table of the Elements

We now understand that atoms and their basic particles are more complicated than the Bohr model, and that protons and neutrons have constituent building parts. Fortunately, this model does describe the properties of elements to a level that explains the properties of different elements.

Elements differ from each other in the number of electrons, protons, and neutrons in their atoms. Nature combines the subatomic particles in an orderly

fashion. An examination of some of the rules governing atomic structure is helpful in understanding the properties of semiconducting materials and process chemicals. Atoms (and therefore the elements) range from the simplest, hydrogen (with 1 electron), to the most complicated one, lawrencium (with 103 electrons).

Hydrogen consists of only one proton in the nucleus and only one electron. This arrangement illustrates the first of the following rules of atomic structure: 1. In each atom, there are an equal number of protons and electrons.

2. Each element contains a specific number of protons, and no two elements have the same number of protons. Hydrogen, for example, has one proton in its nucleus, while the oxygen atom has eight.

This fact leads to the assignment of numbers to each of the elements. Known as the *atomic number*, it is equal to the number of protons (and therefore electrons) in the atom. The basic reference of the elements is the periodic table ([Fig. 2.2](#)). The periodic table has a box for each of the elements, which is identified by two letters. The atomic number is in the upper left-hand corner of the box. Thus, calcium (Ca) has the atomic number 20, so we know immediately that calcium has 20 protons in its nucleus and 20 electrons in its orbital system.

I																	VIII																									
1 H Hydrogen	2 He Helium															3 Li Lithium	4 Be Beryllium																									
5 B Boron	6 C Carbon	7 N Nitrogen	8 O Oxygen	9 F Fluorine	10 Ne Neon											11 Na Sodium	12 Mg Magnesium																									
13 Al Aluminum	14 Si Silicon	15 P Phosphorus	16 S Sulfur	17 Cl Chlorine	18 Ar Argon	19 K Potassium	20 Ca Calcium	21 Sc Scandium	22 Ti Titanium	23 V Vanadium	24 Cr Chromium	25 Mn Manganese	26 Fe Iron	27 Co Cobalt	28 Ni Nickel	29 Cu Copper	30 Zn Zinc	31 Ga Gallium	32 Ge Germanium	33 As Antimony	34 Se Selenium	35 Br Bromine	36 Kr Krypton																			
37 Rb Rubidium	38 Sr Strontium	39 Y Yttrium	40 Zr Zirconium	41 Nb Niobium	42 Mo Molybdenum	43 Tc Technetium	44 Ru Ruthenium	45 Rh Rhodium	46 Pd Palladium	47 Ag Silver	48 Cd Cadmium	49 In Indium	50 Sn Tin	51 Sb Bismuth	52 Te Tellurium	53 I Iodine	54 Xe Xenon	55 Cs Cesium	56 Ba Barium	57 La Lanthanum	58 Ce Cerium	59 Pr Praseodymium	60 Nd Niodymium	61 Pm Promethium	62 Sm Samarium	63 Eu Europium	64 Gd Gadolinium	65 Tb Terbium	66 Dy Dysprosium	67 Ho Holmium	68 Er Erbium	69 Tm Thulium	70 Yb Ytterbium	71 Lu Lutetium								
72 Hf Hafnium	73 Ta Tantalum	74 W Tungsten	75 Re Rhenium	76 Os Osmium	77 Ir Iridium	78 Pt Platinum	79 Au Gold	80 Hg Mercury	81 Tl Thallium	82 Pb Lead	83 Bi Bismuth	84 Po Polonium	85 At Astatine	86 Rn Radon							87 Fr Francium	88 Ra Radium	89 Ac Actinium						90 Th Thorium	91 Pa Protactinium	92 U Uranium	93 Np Neptunium	94 Pu Plutonium	95 Am Americium	96 Cm Curium	97 Bk Berkelium	98 Cf Californium	99 Es Einsteinium	100 Fm Fermium	101 Md Mendelevium	102 No Nobelium	103 Lr Lawrencium

FIGURE 2.2 Periodic table of the elements.

Neutrons are electrically neutral particles that, along with the protons, make up the mass of the nucleus.

[Figure 2.3](#) shows the atomic structure of elements no. 1, hydrogen; no. 3, lithium; and no. 11, sodium. When constructing the diagrams, several rules

were observed in the placement of the electrons in their proper orbits. The rule is that each orbit (n) can hold $2n^2$ electrons. The solution of the math for orbit no. 1 dictates that the first electron orbit can hold only two electrons. This rule forces the third electron of lithium into the second ring. The rule limits the number of electrons in the second ring to 8 and that of the third ring to 18. So, when constructing the diagram of the sodium atom with its 11 protons and electrons, the first two orbits take up 10 electrons, leaving the 11th in the third ring.

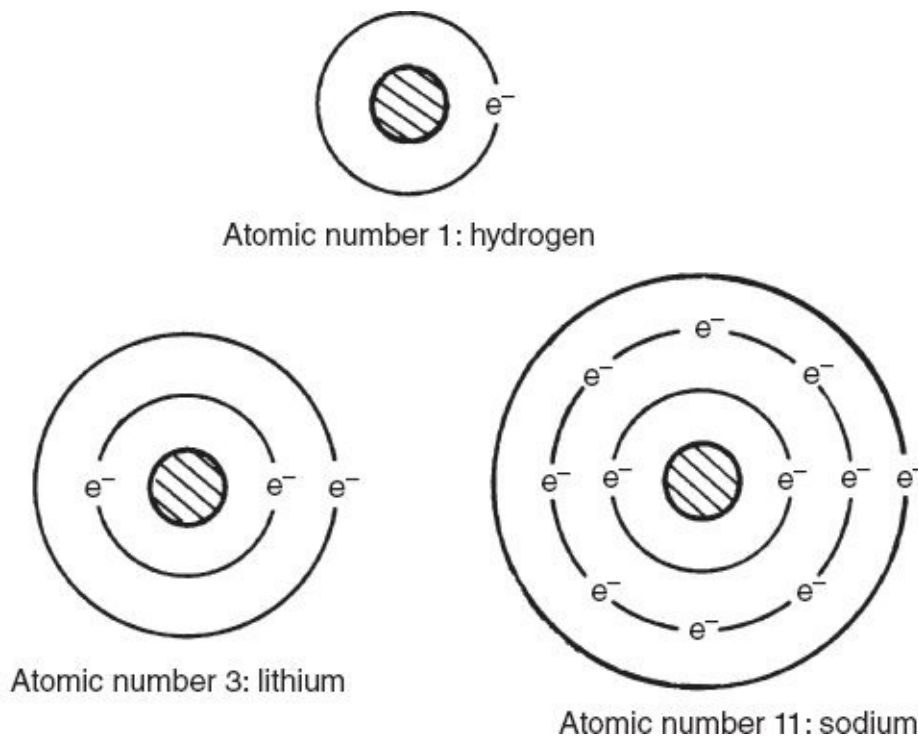


FIGURE 2.3 Atomic structures of hydrogen, lithium, and sodium.

These three atoms have a commonality. Each has an outer ring with only one electron in it. This illustrates another observable fact of elements.

3. Elements with the same number of outer-orbit electrons have similar properties. This rule is reflected in the periodic table. Note that hydrogen, lithium, and sodium appear on the table in a vertical column labeled with the Roman numeral one (I). The column number represents the number of electrons in the outer ring and all of the elements in each column share similar properties.

It is no accident that three of the best electrical conductors (copper, silver, and gold) all appear in the same column (Ib) ([Fig. 2.4](#)) of the periodic table.

There are two more rules of atomic structure that are relevant to the understanding of semiconductors.

29
Cu
Copper
47
Ag
Silver
79
Au
Gold

FIGURE 2.4 The three best electrical conductors.

4. Elements are stable with a filled outer ring or with eight electrons in the outer ring. These atoms tend to be more chemically stable than atoms with partially filled rings.

5. Atoms seek to combine with other atoms to create the stable condition of full orbits or eight electrons in their outer ring.

Rules 4 and 5 influence the creation of N-and P-type semiconductor materials, as explained in the section “Doped Semiconductors.”

Electrical Conduction

Conductors

An important property of many materials is the ability to conduct electricity (i.e., support an electrical current flow). An electrical current is simply a flow of electrons. Electrical conduction takes place in elements and materials where the attractive hold of the protons on the outer ring electrons is relatively weak. In such a material, these electrons can be easily moved, which sets up an electrical current. This condition exists in most metals.

Conductivity is the property of materials to conduct electricity. The higher the conductivity, the better the conductor will be. Conducting ability is also measured by the reciprocal of the conductivity, which is *resistivity*. The lower the resistivity of a material, the better will be the conducting ability of that material.

$$C = 1/\rho$$

where C = conductivity

ρ = resistivity, Ω -cm

Dielectrics and Capacitors

At the opposite end of the conductivity scale are materials that exhibit a large attractive force between the nucleus and the orbiting electrons. The net effect is a resistance to the movement of electrons. These materials are known as *dielectrics*, and in electrical circuits a dielectric functions as an *insulator*. They have low conductivity and high resistivity. In electrical circuits and products, dielectric materials such as silicon dioxide (glass) are used as insulators.

An electrical device known as a *capacitor* is formed whenever a dielectric layer is sandwiched between two conductors. In semiconductor structures, capacitors are formed in MOS gate structures, between metal layers and silicon substrates separated by dielectric layers, and in other structures (see [Chap. 16](#)). The practical effect of a capacitor is that it stores electrical charges. Capacitors are used for information storage in memory devices, to prevent unwanted charges to build up in conductors and silicon surfaces, and to form the working parts of field effect (MOS) transistors. The capacitance ability of a thin dielectric film is relative to the area and thickness and a property parameter known as the *dielectric constant*. Semiconductor metal conduction systems need high conductivity, and therefore low-resistance and low-capacitance materials. These are referred to as *low-k dielectrics*. Dielectric layers used as insulators between conducting layers need high capacitances or *high-k dielectrics*.

$$C = \frac{kE_0A}{t}$$

where C = capacitance

k = dielectric constant of material

E_0 = permittivity of free space (free space has the highest “capacitance”)

A = area of capacitor

t = thickness of dielectric material

Resistors

An electrical factor related to the degree of conductivity (and resistivity) of a material is the electrical resistance of a specific volume of the material. The resistance is a factor of both the resistivity and the dimensions of the material.

Resistance to electrical flow is measured in ohms as illustrated in [Fig. 2.5](#).

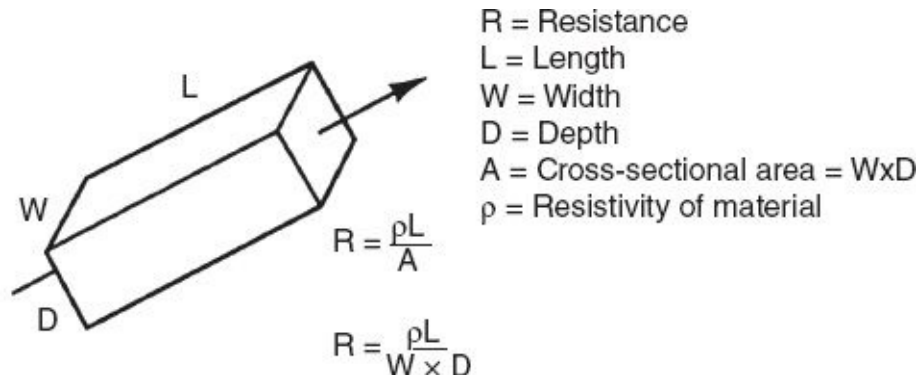


FIGURE 2.5 Resistance of a rectangular bar.

The formula defines the electrical resistance of a specific volume of a specific material (in [Fig. 2.5](#), the volume is a rectangular bar with dimensions W , L , and D). The relationship is analogous to density and weight, with *density* being a material property and *weight* being the force exerted by a specific volume of the material.

Electric current flow is analogous to water flowing in a hose. For a given hose diameter and water pressure, only a given amount of water will flow out of the hose. The resistance to flow can be reduced by increasing the hose diameter, shortening the hose, and/or increasing the pressure. In an electrical system, the electron flow can be increased by increasing the cross-section of the material, shortening the length of the piece, increasing the voltage (analogous to pressure), and/or decreasing the resistivity of the material.

Intrinsic Semiconductors

Semiconducting materials, as the name implies, are materials that have some natural electrical conducting ability. There are two elemental semiconductors (silicon and germanium), and both are found in column IV ([Fig. 2.6](#)) of the periodic table. In addition, there are some tens materials (chemical compounds) that also exhibit semiconducting properties. These compounds come from elements found in columns III and V, such as gallium arsenide (GaAs), indium gallium phosphide (InGaP), and gallium phosphide (GaP). Others are compounds from elements from columns II and VI of the periodic table.

IIIA	IVA	VA
5 10.811 B Boron	6 12.01115 C Carbon	7 14.0067 N Nitrogen
13 26.9815 Al Aluminum	14 28.086 Si Silicon	15 30.9738 P Phosphorus
31 69.72 Ga Gallium	32 72.64 Ge Germanium	33 74.922 As Arsenic
49 114.82 In Indium	50 118.69 Sn Tin	51 121.75 Sb Antimony
81 204.37 Tl Thallium	82 207.19 Pb Lead	83 208.980 Bi Bismuth

Elemental
semiconductors

III - V Compound
semiconductors

FIGURE 2.6 Semiconductor materials.

The term *intrinsic* refers to these materials in their purified state and not contaminated with impurities (dopants) purposely added to change properties.

Doped Semiconductors

Semiconducting materials in their intrinsic state are not useful in solid-state devices. However, through a process called *doping*, specific elements are introduced into intrinsic semiconductor materials. These elements increase the conductivity of the intrinsic semiconductor material. The doped material displays two unique properties that *are* the basis of solid-state electronics. The two properties are: 1. Precise resistivity control through doping

2. Electron and hole conduction

Resistivity of Doped Semiconductors

Metals have a conductivity range limited to 10^4 to 10^6 per ohm-centimeter. The implications of this limit are illustrated by an examination of the resistor represented in [Fig. 2.5](#). Given a specific metal with a specific resistivity, the only way to change the resistance of a given volume is to change the dimensions. In a semiconductive material, the resistivity can be changed, giving another degree of freedom in the design of the resistor. Semiconductors are such materials. Their resistivity can be extended over the range of 10^{-3} to 10^3 by the addition of dopant atoms.

Semiconducting materials can be doped into a useful resistivity range by elements that make the material either electron-rich (N-type) or hole-rich (P-type).

[Figure 2.7](#) shows the relationship of the doping level to the resistivity of silicon. The X-axis is labeled the carrier concentration because the electrons or holes in the material are called *carriers*. Note that there are two curves: N-type and P-type. That is due to the different amount of energies required to move an electron or a hole through the material. As the curves indicate, it takes less of a concentration of N-type dopants than P-type dopants to create a given resistivity in silicon. Another way to express this phenomenon is that it takes less energy to move an electron than to move a hole.

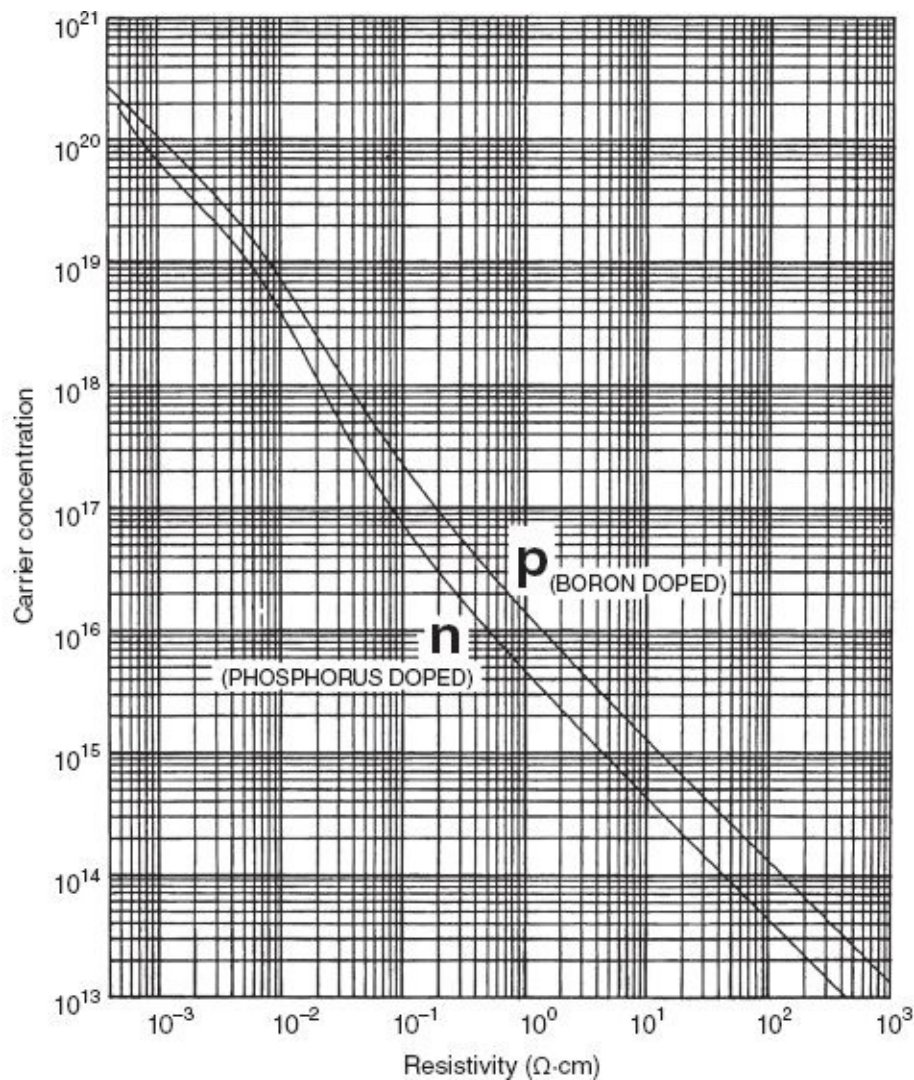


FIGURE 2.7 Silicon resistivity versus doping (carrier) concentration. (After R.L. Thurber et al., Natl. Bur. Std. Spec. Publ., May 1981, Tables 10 and 14: 400–464.)

It takes only 0.000001 to 0.1 percent of a dopant to bring a semiconductor material into a useful resistivity range. This property of semiconductors allows the creation of regions of very precise resistivity values in the material.

Electron and Hole Conduction

Another limit of a metal conductor is that it conducts electricity only through the movement of electrons. Metals are permanently N-type. Semiconductors can be made either N-or P-type by doping with specific dopant elements. N-and P-type semiconductors can conduct electricity by either electrons or holes. Before examining the conduction mechanism, it is instructive to examine the creation of free (or extra) electrons or holes in a semiconductor structure.

To understand the situation of N-type semiconductors, consider a piece of silicon

(Si) doped with a very small amount of arsenic (As), as shown in [Fig. 2.8](#). Assuming even mixing, each of the arsenic atoms would be surrounded by silicon atoms. Applying the rule from the “Periodic Table of the Elements” section that atoms attempt to stabilize by having eight electrons in their outer ring, the atom is shown sharing four electrons from its neighboring silicon atoms. However, arsenic is from column V, which means it has five electrons in its outer ring. The net result is that four of them pair up with electrons from the silicon atoms, leaving one left over. This one electron is available for electrical conduction.

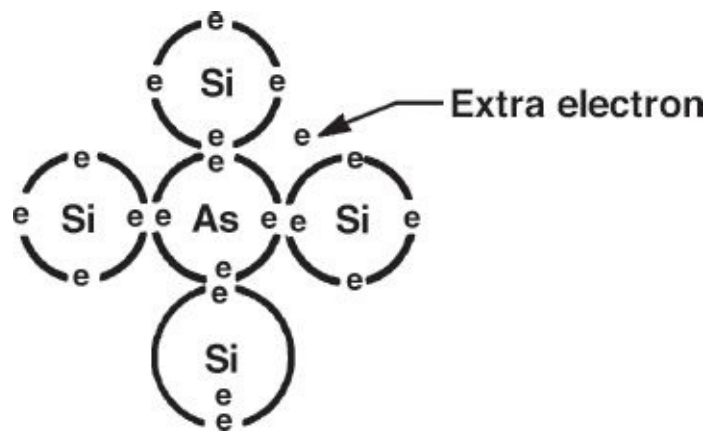


FIGURE 2.8 N-type doping of silicon with arsenic.

Considering that a crystal of silicon has millions of atoms per cubic centimeter, there are lots of electrons available to conduct an electrical current. In silicon, the elements arsenic, phosphorus, and antimony create N-type conditions.

An understanding of P-type material is approached in the same manner ([Fig. 2.9](#)). The difference is that only boron, from column III of the periodic table, is used to make silicon P-type. When mixed into silicon, boron too borrows electrons from silicon atoms. However, having only three outer electrons, there is a place in the outer ring that is not filled by an electron. This unfilled position is defined as a *hole*.

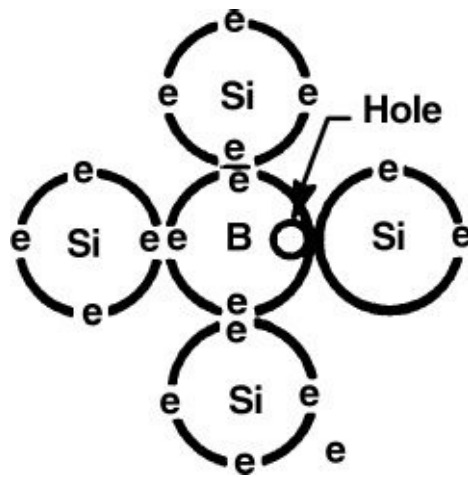


FIGURE 2.9 P-type doping of silicon with boron.

Within a doped semiconductor material, there is a great deal of activity—holes and electrons are constantly being created. The electrons are attracted to the unfilled holes, in turn leaving an unfilled position, which creates another hole.

How the electrons contribute to electrical conduction is illustrated in [Fig. 2.10](#). When a voltage is applied across a piece of conducting or semiconducting material, the negative electrons move toward the positive pole of the voltage source, such as a battery.

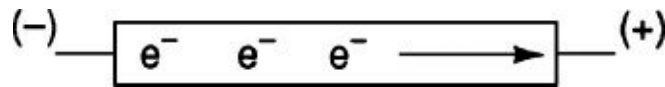


FIGURE 2.10 Electron conduction in N-type semiconductor material.

In P-type material ([Fig. 2.11](#)), an electron will move toward the positive pole by jumping into a hole along the direction of the route (t_1).

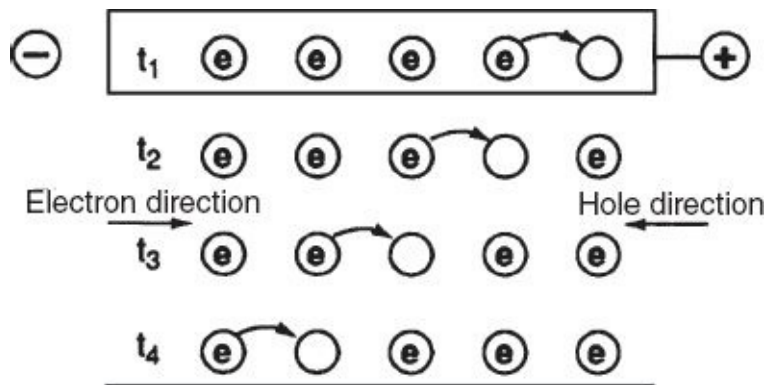


FIGURE 2.11 Hole conduction in P-type semiconductor material.

Of course, when an electron leaves its position, it leaves a new hole. As it

continues toward the positive pole, it creates a succession of holes. The effect to someone measuring this process with a current meter is that the material is supporting a positive current, when actually it is a negative current moving in the opposite direction. This phenomenon is called a *hole flow* and is unique to semiconducting materials.

The dopants that create a P-type conductivity in a semiconductor material are called *acceptors*. Dopants that create N-type conditions are called *donors*. An easy way to keep these terms straight is that acceptor has a p and donor is spelled with an n.

The electrical characteristics of conductors, insulators, and semiconductors are summarized in [Fig. 2.12](#). The particular characteristics of doped semiconductors are summarized in [Fig. 2.13](#).

Classification	Electrons	Examples	Conductivity
1. Conductor	Free to Move	Gold Copper Silver	$10^4 - 10^6$ /ohm-cm
2. Insulator (Dielectric)	Bound	Glass Plastic	$10^{-22} - 10^{-10}$ /ohm-cm
3. Semiconductor a. Intrinsic	Some Available	Germanium Silicon III-IV	$10^{-9} - 10^3$ /ohm-cm
b. Doped	Controlled Amount Available	N-type Semiconductor P-type Semiconductor	

FIGURE 2.12 Electrical classification of materials.

	N-TYPE	P-TYPE
1. Conduction	Electrons	Holes
2. Polarity	Negative	Positive
3. Dopant Term	Donor	Acceptor
4. Doping Elements in Silicon	Arsenic Phosphorus Antimony	Boron

FIGURE 2.13 Characteristics of doped semiconductors.

N-and P-type conditions are also created with specific elements, in germanium and compound semiconductors.

Carrier Mobility

It was mentioned previously that less energy is needed to move an electron than a hole through a piece of semiconducting material. In a circuit we are interested in both the energy required to move these carriers (holes and electrons), and the speed at which they move. The speed of movement is called the *carrier mobility*, with holes having a lower mobility than electrons. This factor is an important consideration in selecting a particular semiconducting material for a circuit.

Semiconductor Production Materials

Germanium and Silicon

Germanium and silicon are the two elemental semiconductors. The first transistor was made with germanium, as were the initial devices of the solid-state era. However, germanium presents significant problems in processing and in device performance. For processing, high temperatures are needed. Germanium's 937°C melting point limits it to much lower-temperature processing than silicon. More importantly, its lack of a naturally occurring oxide leaves the surface of a germanium-based chip prone to electrical leakage.

The development of silicon/silicon dioxide planar processing solved the leakage problem of integrated circuits, flattened the surface profile of the circuits, and allowed higher temperature processing due to its 1415°C melting point. Consequently, silicon represents over 90 percent of the wafers processed worldwide.

Semiconducting Compounds

There are many semiconducting compounds formed by combining elements from columns III and V, and II and VI of the periodic table. Of these compounds, the ones most used in commercial semiconductor devices are gallium arsenide (GaAs), gallium arsenide-phosphide (GaAsP), indium phosphide (InP), gallium aluminum arsenic (GaAlAs), and indium gallium phosphide (InGaP).¹ These compounds have special properties.² For example, diodes made from GaAs and GaAsP give off visible and laser light when activated with an electrical current. They and other materials are used to make the light-emitting diodes (LEDs). LEDs have a growing market since the development of additional compounds that give off a range of colored light. An annotated list is in [Fig 2.14](#).

	Ge	Si	GaAs	SiO ₂
Atomic Weight	72.6	28.09	144.63	60.08
Atoms/cm ³ or Molecules	4.42×10^{22}	5.00×10^{22}	2.21×10^{22}	2.3×10^{22}
Crystal Structure	Diamond	Diamond	Zinc-Blends	Amorphous
Atoms/Unit Cell	8	8	8	—
Density	5.32	2.33	5.65	2.27
Energy Gap	0.67	1.11	1.40	8 (approx.)
Dielectric Constant	16.3	11.7	12.0	3.9
Melting Point (°C)	937°	1415°	1238°	1700° (approx.)
Breakdown Field (V/cm)	8 (approx.)	30 (approx.)	35 (approx.)	600 (approx.)
Linear Coefficient of Thermal Expansion	5.8×10^{-6}	2.5×10^{-6}	5.9×10^{-6}	0.5×10^{-6}
$\frac{\Delta L}{L T}$	$\frac{1}{C}$			

FIGURE 2.14 Physical properties of semiconductor materials.

An important property of gallium arsenide is its high (electrical) carrier mobility. This property allows a gallium arsenide device to react to high-frequency microwaves and effectively switch them into electrical currents in communications systems faster than can be done in silicon devices.

This same property, carrier mobility, is the basis for the excitement over gallium arsenide transistors and ICs. Devices of GaAs operate two to three times faster than comparable silicon devices and find applications in super-fast computers and real-time control circuits such as airplane controls.

GaAs has a natural resistance to radiation-caused leakage. Radiation, such as that found in space, causes holes and electrons to form in semiconductor materials. It gives rise to unwanted currents that can cause the device or circuit to malfunction or cease functioning. Devices that can perform in a radiation environment are known as *radiation hardened*. GaAs is naturally radiation hardened.

GaAs is also semi-insulating. In an integrated circuit, this property minimizes leakage between adjacent devices, allowing a higher packing density. Higher packing density in turn results in a faster circuit because the holes and electrons travel shorter distances. In silicon circuits, special isolating structures must be built into the surface to control surface leakage. These structures take up valuable space and reduce the density of the circuit.

Despite all of the advantages, GaAs is not expected to replace silicon as the mainstream semiconducting material. The reasons reside in the trade-offs between performance and processing difficulty. While GaAs circuits are very fast, the majority of electronic products do not require their level of speed. On the performance side, GaAs, like germanium, does not possess a natural oxide. To compensate, layers of dielectrics must be deposited on the GaAs, which leads to longer processing and lower yields. Also, half of the atoms in GaAs are arsenic, an element that is very dangerous to human beings. Unfortunately, the arsenic evaporates from the compound at normal process temperatures, thus requiring for human safety an addition of suppression layers (caps) or pressurized process chambers. These steps lengthen the processing and add to its cost.

Evaporation also occurs during the crystal-growing stage of GaAs, resulting in nonuniform crystals and wafers. The nonuniformity produces wafers that are very prone to breakage during fab processing. Also, the production of large-diameter GaAs wafers has lagged behind that of silicon (see [Chap. 3](#)).

Despite the problems, gallium arsenide is an important semiconducting material that will continue to increase in use and will probably have a major influence on computer performance of the future.

Silicon Germanium

Competitors to GaAs are silicon-germanium (SiGe) structures. The combination increases transistor speeds to levels that allow ultra-fast radios and personal communication devices.³ Device/IC structures feature a layer of germanium deposited by ultrahigh vacuum/chemical vapor deposition (UHV/CVD).⁴ Bipolar transistors are formed in the Ge layer. Unlike the simpler transistors formed in silicon technology, SiGe requires transistors with *heterostructures* or *heterojunctions*. These are structures with several layers and specific dopant levels to allow high-frequency operations.

A comparison of the major semiconducting production materials and silicon dioxide is presented in [Fig. 2.14](#).

Engineered Substrates

A bulk wafer was the traditional substrate for fabricating microchips for many years. Electrical performance now often demands new substrates, such as silicon on an insulator (SOI) such as sapphire, and silicon on diamond (SOD). Diamond dissipates heat better than silicon. Another structure is a layer of *strained* silicon deposited on a wafer of silicon germanium. Strained silicon occurs when silicon atoms are deposited on an Si/Ge (SOI) layer previously deposited on an insulator.

Si/Ge atoms are more widely spaced than normal silicon. During the deposition, the silicon atoms stretch to align to the Si/Ge atoms, straining the silicon layer. The electrical effect is to lower the silicon resistance, allowing electrons to move up to 70 percent faster. This structure brings performance benefits to MOS transistors (see [Chap. 16](#)).

Ferroelectric Materials

In the ongoing search for faster and more reliable memory structures, ferroelectrics have emerged as a viable option. A memory cell must store information in one of two states (on/off, high/low, 0/1), be able to respond quickly (read and write), and be capable of changing states reliably. Ferroelectric material capacitors such as $\text{PbZr}_{1-x}\text{T}_x\text{O}_3$ (PZT) and $\text{SrBi}_2\text{Ta}_2\text{O}_9$ (SBT) exhibit these desirable characteristics. They are incorporated into SiCMOS (see [Chap. 16](#)) memory circuits known as ferroelectric random access memories (FeRAMs).⁵

Diamond Semiconductors

Moore's law cannot go indefinitely into the future. One end point is when the transistor parts become so tiny that the physics governing transistor action no longer work. Another limitation is heat dissipation. Bigger and denser chips run very hot. Unfortunately, high heat degrades the electrical operations and can render the chip useless. Diamond is a crystal material that dissipates heat much faster than silicon. Despite this positive aspect, diamond as a semiconductor wafer has faced barriers of cost, uniformity, and finding a supply of large diamonds. However, there is new research into making less-costly synthetic diamonds using vapor-deposition techniques. Also, research into the doping of diamond into *n*- and *p*-type conductivities is bringing the possibility of diamond semiconductors into reality. This material is being explored and may find its way into fabrication areas of the future.⁶

Process Chemicals

It should be fairly obvious that extensive processing is required to change raw semiconducting materials into useful devices. The majority of these processes use chemicals. In fact, microchip fabrication is primarily a chemical process—or more correctly, a series of chemical processes. Up to 20 percent of all process steps are cleaning or wafer surface preparation.⁷

The cost of processing microchips is getting higher due in part to all the chemicals involved. Great quantities of acids, bases, solvents, and water are

consumed by a semiconductor plant. Part of this is due to the extremely high purities and special formulations required of the chemicals to allow precise and clean processing. Larger wafers and higher cleanliness requirements need more automated cleaning stations. Also, the cost of removal of spent chemicals is rising. When the costs of producing a chip are added up, processing chemicals can be up to 40 percent of all manufacturing costs.

The cleanliness requirements for semiconductor process chemicals are explored in [Chap. 4](#). Specific chemicals and their properties are detailed for specific processes in [Chaps. 7–13](#).

Molecules, Compounds, and Mixtures

At the beginning of this chapter, the basic structure of matter was explained by the use of the Bohr atomic model. This model was used to explain the structural differences of the elements that make up all the materials in the physical universe. But it is obvious that the universe contains more than 103 (the number of elements) types of matter.

The basic unit of a nonelemental material is the molecule. Water, for example, is a molecule composed of two hydrogen atoms and one oxygen atom. The multiplicity of materials in our world comes about from the ability of atoms to bond together to form molecules.

It is inconvenient to draw diagrams such as in [Fig. 2.15](#) every time we want to designate a molecule. The more common practice is to write the molecular formula. For water, it is the familiar H_2O . This formula tells us exactly the elements and their number in the material. Chemists use the more precise term *compound* in describing different combinations of elements. Thus, H_2O (water), NaCl (sodium chloride or salt), H_2O_2 (hydrogen peroxide), and As_2O_3 (arsine) are all different compounds composed of aggregates of individual molecules.

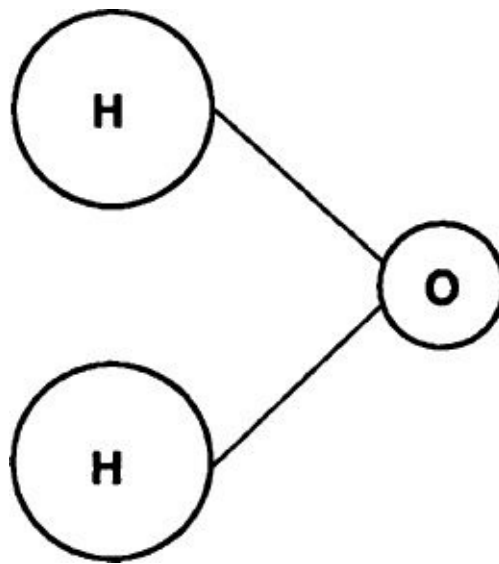


FIGURE 2.15 Diagram of a water molecule.

Some elements combine into *diatomic molecules*. A diatomic molecule is one composed of two atoms of the same element. The familiar process gases (oxygen, nitrogen, and hydrogen), in their natural state, are all composed of diatomic molecules. Thus, their formulas are O₂, N₂, and H₂.

Materials also come in two other forms: mixtures and solutions. Mixtures are composed of two or more substances, but the substances retain their individual properties. A mixture of salt and pepper is a classic example.

Solutions are mixtures of a solid dissolved in a liquid. In the liquid, the solids are interspersed, with the solution taking on unique properties. However, the substances in a solution do not form a new molecule. Saltwater is an example of a solution. It can be separated back into its starting parts: salt and water.

Slurries are considered a subtype of mixture. Slurries combine a solid with a liquid. The solid does not dissolve, and the individual materials each retain their individual properties. Slurries are used in polishing operations such as *chemical mechanical polishing* (CMP). Typical processing slurries have fine pieces of silica (glass) suspended in a mild base solution such as ammonium hydroxide.

Ions

The term *ion* or *ionic* is used often in connection with semiconductor processing. This term refers to any atom or molecule that exists in a material with an unbalanced charge. An ion is designated by the chemical symbol of the element or molecule, followed by a superscripted positive or negative sign (e.g., Na⁺, Cl⁻). For example, one of the serious contamination problems is mobile ionic contamination such as sodium (Na⁺). The problem comes from the positive charge carried by the sodium when it gets into the semiconductor material or device. In some processes,

such as the ion-implantation process, it is necessary to create an ion, such as boron (B^+), to accomplish the process.

States of Matter

Solids, Liquids, and Gases

Matter is found in four different states. They are solids, liquids, gases, and plasma ([Fig. 2.16](#)).

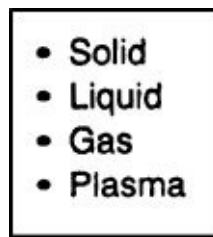


FIGURE 2.16 Four states of nature.

- Solids are defined as having a definite shape and volume, which is retained under normal conditions of temperature and pressure.
- Liquids have definite volume but a variable shape. A liter of water, for example, will take the shape of any container in which it is stored.
- Gases have neither a definite shape, nor volume. They too will take the shape of any container but, unlike liquids, they will expand or can be compressed to entirely fill the container.

The state of a particular material has a lot to do with its pressure and temperature. Temperature is a measure of the total energy incorporated in the material. We know that water can exist in three states (ice, liquid water, and steam or water vapor) simply by changing the temperature and/or pressure. The influence of pressure is very complicated and beyond the scope of this text.

Plasma State

The fourth state of nature is plasma. A star is an example of a plasma state. It certainly does not meet the definitions of a solid, liquid, or gas. A plasma state is defined as a high-energy collection of ionized atoms or molecules. Plasma states can be induced in process gases by the application of high-energy radio-frequency (RF) fields. They are used in semiconductor technology to cause chemical reactions in gas mixtures. One of their advantages is that energy can be delivered at a lower

temperature than in convention systems, such as convection heating in ovens.

Properties of Matter

All materials can be differentiated by their chemical compositions and the properties that arise from those compositions. In this section, several key properties are defined, that are required to understand and work with properties of semiconductor materials and chemicals.

Temperature

The temperature of a chemical exerts great influence on that chemical's reactions with other chemicals, whether in an oxidation tube or in a plasma etcher. Additionally, safe use of some chemicals requires knowledge and control of their temperatures. Three temperature scales are used to express the temperature of a material. They are the Fahrenheit, Centigrade (or Celsius), and the Kelvin scales ([Fig. 2.17](#)).

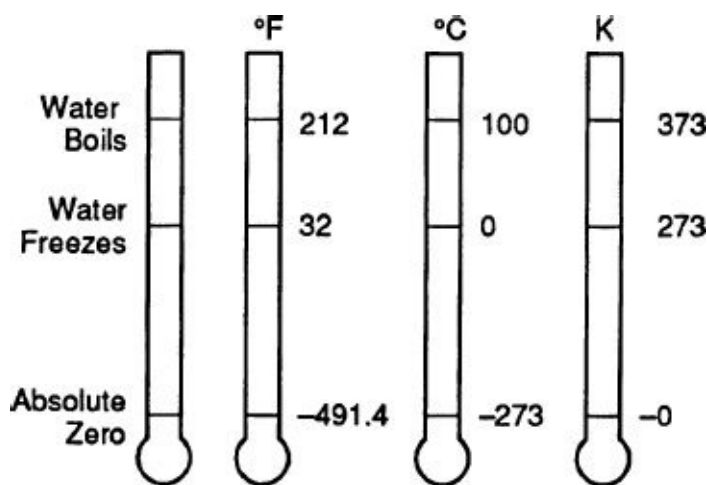


FIGURE 2.17 Temperature scales.

The Fahrenheit scale was developed by Gabriel Fahrenheit, a German physicist, using a water and salt solution. He assigned to the solution's freezing temperature the value of zero degrees Fahrenheit (0°F). Unfortunately, the freezing temperature of pure water is more useful, and we have ended up with the Fahrenheit scale having a water freezing point at 32°F and a boiling point of 212°F , with 180° between the two points.

The Celsius or centigrade scale is more popular in scientific endeavors. It more sensibly sets the freezing point of water at 0°C and boiling at 100°C . Note that there are exactly 100°C between the two points. This means that a 1° change in

temperature as measured on the centigrade scale requires more energy than a 0° change on the Fahrenheit scale.

The third temperature scale is the Kelvin scale. It uses the same scale factor as the centigrade scale but is based on absolute zero. Absolute zero is the theoretical temperature at which all atomic motion would cease. This value corresponds to – 273°C. On the Kelvin scale, water freezes at 273 K and boils at 373 K.

Density, Specific Gravity, and Vapor Density

An important property of matter is density. When we say that something is *dense*, we refer to its mass or weight per unit volume. A cork has lower density than an equal volume of iron. Density is expressed as the weight in grams, per cubic centimeter of the material. Water is the standard, with 1 cm³ (at 4°C) weighing 1 g. The densities of other substances are expressed as a ratio of their density to that of a comparable volume of water. Silicon has a density of 2.3. Therefore, a piece of silicon 1 cm³ (one cubic centimeter) in volume will weigh 2.3 g.

Specific gravity is a term used to reference the density of liquids and gases at 4°C. It is the ratio of the density of a substance compared to that of water. Gasoline has a specific gravity of 0.75, which means it is 75 percent as dense as water.

Vapor density is a density measurement of gases under certain conditions of temperature and pressure. The reference is air, with one 1 cm³ having an assigned density of 1 (one). Hydrogen has a vapor density of 0.60 which makes it 60 percent the density of a similar volume of air. The contents of a container of hydrogen will therefore weigh 60 percent less than a similar container of air.

Pressure and Vacuum

Another important aspect of matter is pressure. Pressure, as a property, is usually used in connection with liquids and gases. It is defined as the force per unit area exerted against the surface of the container. It is the gas pressure in a cylinder that forces the gas out into a process chamber. All process machines using gases must have gauges to measure and control the pressure.

Pressures are expressed in pounds per square inch of area (psia), in atmospheres or in torrs. An atmosphere (atm) is the pressure exerted by the atmosphere surrounding the Earth at a specific temperature. Thus, a high-pressure oxidation system operated at 5 atm contains a pressure 5 times that of the atmosphere.

One atmosphere of air has a pressure of 14.7 psia. Pressures inside gas tanks are measured in psig units, or pounds per square inch gauge. This means that the gauge reading is absolute; it does not include the pressure of the outside atmosphere.

Vacuum is also a term and condition encountered in semiconductor processing. It

is actually a condition of low pressure. Generally, pressures below standard atmospheric pressures are referred to as vacuums. But a vacuum condition is measured in units of pressure.

Low pressures tend to be expressed in torrs. The unit is named after the Italian scientist, Torricelli, who made many of the early discoveries in the field of gases and their properties. A torr is defined as the equivalent of 1 mm of mercury in a pressure measuring device known as a *manometer*.

Imagine the effect on the column of mercury in the manometer in [Fig. 2.18a](#) of increasing the pressure above atmospheric pressure. As the pressure goes up, it pushes down the mercury in the dish and raises the mercury in the column. Now imagine what happens as air is extracted from the system ([Fig. 2.18b](#)) below atmospheric pressure, creating a vacuum. As long as there are any gas molecules or atoms in the manometer, some small pressure will be exerted on the mercury in the dish, and the mercury in the column will rise some small but finite amount. The amount of the rise as measured in millimeters (mm) is relative to the pressure or, in this case, the vacuum.

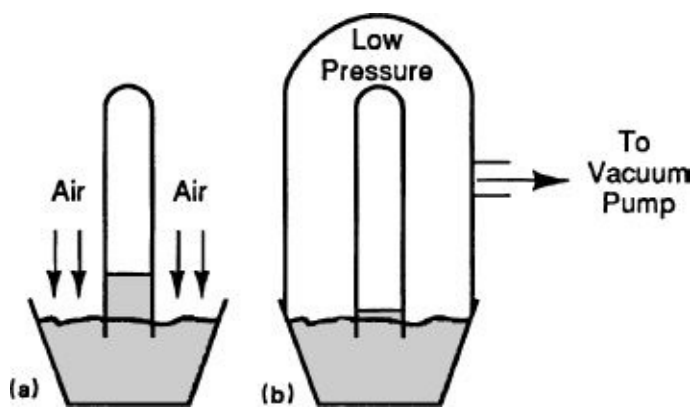


FIGURE 2.18 Pressure vacuum measurement.

Vacuum systems for evaporation, sputtering, and ion implantation are operated at vacuums (pressures) of 10^{-6} to 10^{-9} torr. Translated into a vacuum system containing a simple manometer, this means that the column of mercury would rise only 0.000000001 (1×10^{-9}) to 0.000001 (1×10^{-6}) mm—a *very* short length. In actual practice, a mercury manometer cannot measure these extremely low pressures. Other, more sensitive gauges are used.

Acids, Alkalis, and Solvents

Acids and Alkalis

Semiconductor processing requires large amounts of liquid chemicals to etch, clean, and rinse the wafers and packages. These chemicals are divided by chemists into three major classifications:

- Acids

- Alkalies
- Solvents

Acids and alkalies differ from each other due to the presence of specific ions in the liquid. Acids contain *hydrogen ions*, while alkalies (also called *bases*) contain *hydroxide ions*. An examination of the water molecule explains the differences.

The chemical formula for water normally is written as H_2O . It can also be written in the form HOH . When separated into its parts, we find that water is made up of a positively charged hydrogen ion (H^+) and a negatively charged hydroxide ion (OH^-).

When water is mixed with other elements, either the hydrogen or hydroxyl ion combines with other substances ([Fig. 2.19](#)). Liquids that contain the hydrogen ion are called *acids*. Liquids that contain the hydroxyl ion are called *alkalies* or *bases*. Acids and bases are commonly found in the home: lemon juice and vinegar are acids, and ammonia and baking soda in a solution of water are bases.

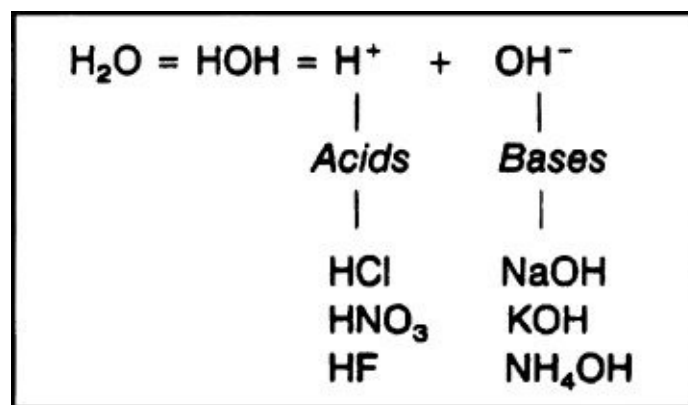


FIGURE 2.19 Acid and base solutions.

Acids are further divided into two categories: organic and inorganic. Organic acids are those that contain hydrocarbons, whereas inorganic acids do not. Sulfonic acid is an organic acid, and hydrogen fluoride (HF) is an inorganic acid.

The strength and reactivity of acids and bases are measured by the pH scale ([Fig. 2.20](#)). This scale ranges from 0 to 14, with 7 being a neutral point. Water is neutral, neither an acid nor a base; therefore, it has a pH of 7. Strong acids, such as sulfuric acid (H_2SO_4), will have low pH values of 0 to 3. Strong bases, such as sodium hydroxide (NaOH), have pH values greater than 7.

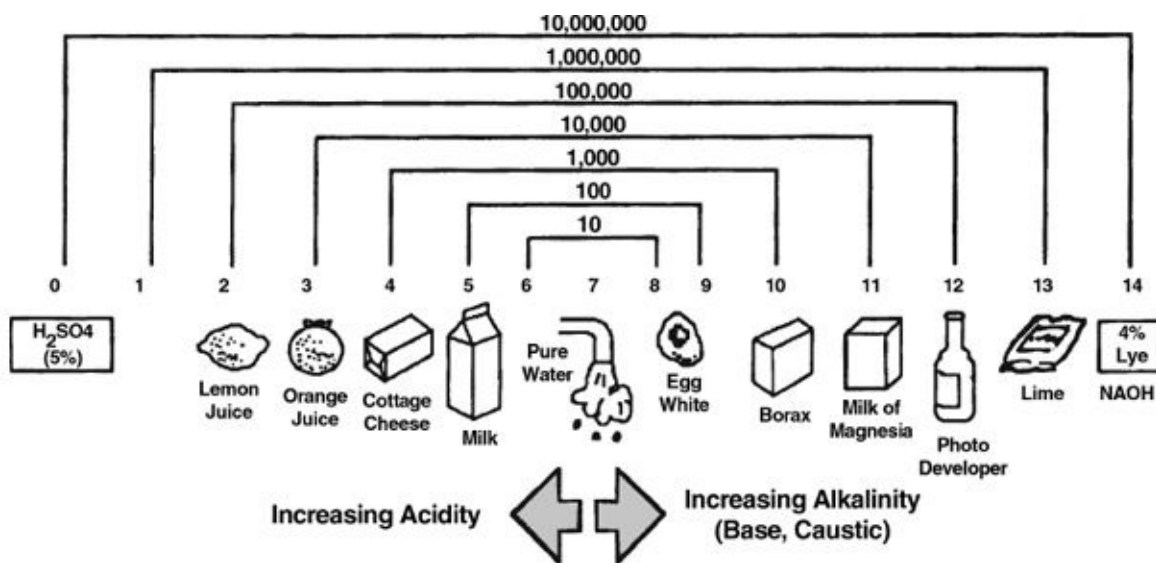


FIGURE 2.20 pH scale.

Both acids and bases are reactive with skin and other chemicals and should be stored and handled with all of the prescribed safety precautions.

Solvents

Solvents are liquids that do not ionize; they are neutral on the pH scale. Water is a solvent; in fact, it is the solvent with the greatest ability to dissolve other substances. It is also the most commonly used solvent in semiconductor processing. Alcohol and acetone are other common solvents in the wafer-fabrication process.

Most of the solvents in fab processing are volatile, flammable, or combustible. It is important to use them in properly exhausted stations and observe prescribed precautions in their storage and use.

Chemical Purity and Cleanliness

While the names of chemicals used in fabrication areas sound familiar, there is an entire supply industry dedicated to producing the highest quality chemicals to meet semiconductor processing demands. Wafers getting up to 100 cleans in fabrication and contamination problems are getting harder to detect.⁸ Chemicals must meet very high purity requirements both chemically and physically. In general the target is “six nines” purity. This translates to 99.9999 percent pure. Physical contamination such as unwanted particles is also controlled. Typical chemical specifications limit particles to parts per billion (ppb)/liter and microns. These and other specifications are established in the *International Technology Roadmap for Semiconductors* (ITRS). Specific chemicals used are identified in the chapter on contamination control, and in the process [Chaps. 7–13](#).

Safety Issues

The storage, use, and disposal of chemicals, and electrical and other risks are present in semiconductor process areas. Companies address these risks by developing employee knowledge, skill, and awareness through training programs and safety inspections.

The Material Safety Data Sheet

For every chemical brought into a manufacturing site, the supplier must provide a form called the Material Safety Data Sheet (MSDS). It is required by the federal Occupational, Safety, and Health Administration (OSHA). The form is also called OSHA Form 20. Each form contains chemical, storage, health, first aid, and usage information about a specific chemical. Under current regulations, the MSDSs must be filed on the site and available to employees.

Review Topics

Upon completion of this chapter, you should be able to:

1. Identify the parts of an atom.
2. Name the two unique properties of a doped semiconductor.
3. List at least three semiconducting materials.
4. Explain the advantages and disadvantages of gallium arsenide compared with silicon.
5. Explain the difference in composition and electrical functioning of N- and P-type semiconducting materials.
6. Describe the properties of resistivity and resistance.
7. Identify the differences between acids, alkalis, and solvents.
8. List the four states of nature.
9. Give the definition of an atom, a molecule, and an ion.
10. Explain four or more basic chemical handling safety rules.

References

1. Fujitsu Quantum Devices Limited, www.datasheetarchive.com/Fujitsu+Quantum+Devices-datasheet.html, 2004.
2. Williams, R. E., *Gallium Arsenide Processing Techniques*, Artech House, Inc., Dedham, MA, 1984.

[3.](#) Ouellette, J., "Silicon–Germanium Gives Semiconductors the Edge," *The Industrial Physicist*, June/July, 2002.

[4.](#) Holton, W. C., "Silicon Germanium: Finally for Real," *Solid State Electronics*, Nov. 1997:119.

[5.](#) R., E., "Integration of Ferroelectrics into Nonvolatile Memories," *Solid State Technology*, Oct. 1997:201.

[6.](#) Smith, J. E., "81 GHz Diamond Semiconductor Created," *Geek.com* (accessed: August 27, 2003).

[7.](#) Allen, R., "MNST Wafer Cleaning," *Solid State Technology*, Jan. 1994:61.

[8.](#) Ibid.

CHAPTER 3

Crystal Growth and Silicon Wafer Preparation

Introduction

In this chapter, the preparation of semiconductor-grade silicon from sand, its conversion into crystals and wafers (material preparation stage), and the processes required to produce polished wafers (crystal growth and wafer preparation) are explained. Included are descriptions of the different types of wafers used in a fabrication operation. The challenges of growing 450-mm diameter crystals and preparing 450-mm wafers are presented.

The evolution of higher-density and larger-size chips has required the delivery of larger diameter wafers. Starting with 1-in diameter wafers in the 1960s, the industry moved to 300-mm (12-in) diameter wafers in the 2000s and is now moving into the 450-mm (18-in) world ([Fig. 3.1](#)).

Year of Production	2001	2006	2012	2015
Wafer Diameter (mm)	300	300	450	450

FIGURE 3.1 Wafer diameters. (Courtesy of SIA.)

Larger-diameter wafers are necessary to accommodate increasing chip sizes with cost-effective wafer-fabrication processes (see [Chaps. 6](#) and [15](#)). The challenges in wafer preparation are formidable. In crystal growth, the issues of structural and electrical uniformity and contamination are challenges. In wafer preparation, flatness, diameter control, impurity content, and crystal integrity are issues. Larger diameters are heavier, which requires more substantial process tools and, ultimately, full automation. A production lot of 300-mm diameter wafers weighs about 20 lb (7.5 kg) and can be worth half a million dollars or more.¹ A 450-mm weighs around 800 kg and is 210 cm long.² These challenges coexist with ever-tightening specifications for almost every parameter. Keeping abreast of these challenges as diameters have grown is a key to continued microchip evolution. However, the conversion to larger diameter wafers is costly and time-consuming. Thus some fabs are still using a smaller range of wafer diameter ([Fig. 3.2](#)) as the industry phases into larger diameter wafers.

Diameter (mm)	Diameter (inch)	Thickness (um)
150 mm	5.9 inch (6 inch)	625 um
200 mm	7.9 inch (8 inch)	725 um
300 mm	11.8 inch (12 inch)	775 um
450 mm	17.7 inch (18 inch)	925 um-projected

FIGURE 3.2 Table of wafer diameters and thickness.

Semiconductor Silicon Preparation

Semiconductor devices and circuits are formed in and on the surface of wafers of a semiconductor material, usually silicon. Those wafers must have very low levels of contaminants, be doped to a specified resistivity level, have a specific crystal structure, be optically flat, and meet a host of other mechanical and cleanliness specifications.

Silicon Wafer Preparation Stages

Manufacture of IC grade silicon wafers proceeds in four stages: • Conversion of ore to a high-purity gas • Conversion of gas to polysilicon silicon • Conversion of polysilicon silicon to a single crystalline, doped crystal ingot • Preparation of wafers from the crystal ingot The first stage of semiconductor manufacturing is the extraction and purification of the raw semiconductor material(s) from the earth. Purification starts with a chemical reaction. For silicon, it is the conversion of the ore to a silicon-bearing gas such as silicon tetrachloride or trichlorosilane. The silicon-bearing gas is then reacted with hydrogen ([Fig. 3.3](#)) to produce semiconductor-grade silicon. The silicon produced is 99.9999999 percent pure; one of the purest materials on Earth.³ Its crystal structure is known as polycrystalline or *polysilicon*.

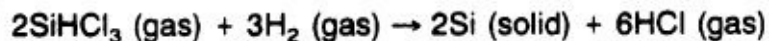


FIGURE 3.3 Hydrogen reduction of trichlorosilane.

Crystalline Materials

One way that materials differ is in the organization of their atoms. In some materials, such as silicon and germanium, the atoms are arranged into a very definite structure that repeats throughout the material. These materials are called

crystals.

Materials without a definite periodic arrangement of their atoms are called noncrystalline or *amorphous*. Plastics are examples of amorphous materials.

Unit Cells

Two levels of atomic organization are possible for crystalline materials. First is the *unit cell* in which the atoms are arranged at specific points in a specific shape. Another term used to reference crystal structures is *lattice*. A crystalline material is said to have a specific lattice structure and the atoms are located at specific points in the lattice structure. The number of atoms, relative positions, and binding energies between the atoms in the unit cell give rise to many of the characteristics of the material. Each crystalline material has a unique unit cell. Silicon atoms have 16 atoms arranged into a diamond structure ([Fig. 3.4](#)). GaAs crystals have 18 atoms in a unit cell configuration called a zincblend ([Fig. 3.5](#)).

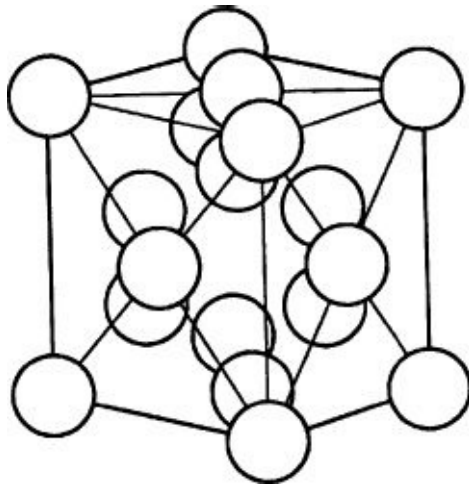


FIGURE 3.4 Unit cell of silicon.

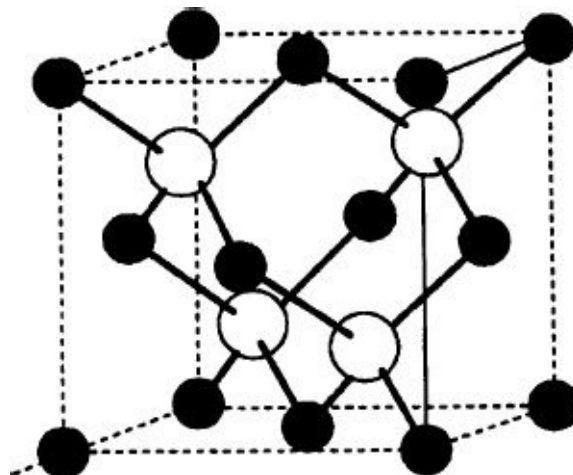


FIGURE 3.5 GaAs crystal structure.

Poly and Single Crystals

The second level of organization within a crystal is related to the organization of the unit cells. In intrinsic semiconductors, the unit cells are *not* in a regular arrangement with each other. The situation is similar to a disorderly pile of sugar cubes, with each cube representing a unit cell. A material with such an arrangement has a polycrystalline structure.

The second level of organization occurs when the unit cells (sugar cubes) are all neatly and regularly arranged relative to each of the others ([Fig. 3.6](#)). Materials thus arranged have a single- (or mono-) crystalline structure.

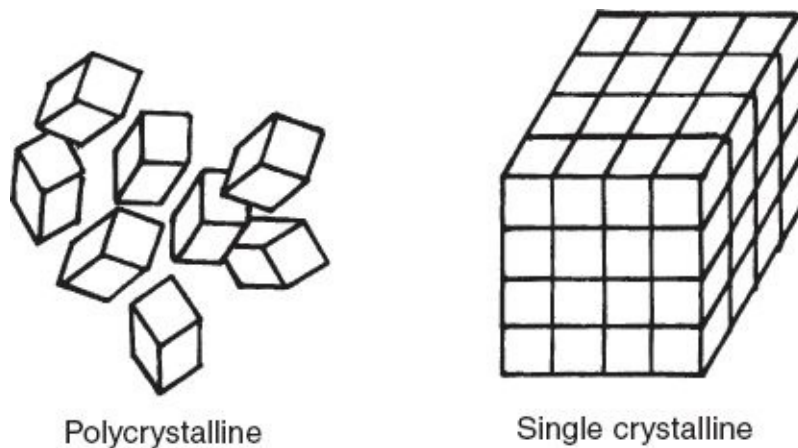


FIGURE 3.6 Poly- and single-crystal structures.

Single-crystal materials have more uniform and predictable properties than polycrystalline materials. The structure allows a uniform and predictable electron flow in semiconductors. At the end of the fab process, crystal uniformity is essential for separating the wafer into die with clean edges and dimensionally correct sides (see [Chap. 18](#)).

Crystal Orientation

In addition to the requirement of a single-crystal structure for a wafer, there is the requirement of a specific *crystal orientation*. This concept can be visualized by considering slicing the single-crystalline block as shown in [Fig. 3.7](#). Slicing it in the vertical plane would expose one set of *planes*, while slicing it corner-to-corner would expose a different plane. Each plane is unique, differing in atom count and binding energies between the atoms. Each has different chemical, electrical, and physical properties that are imparted to the wafers. Specific crystal orientations are required for the wafers.

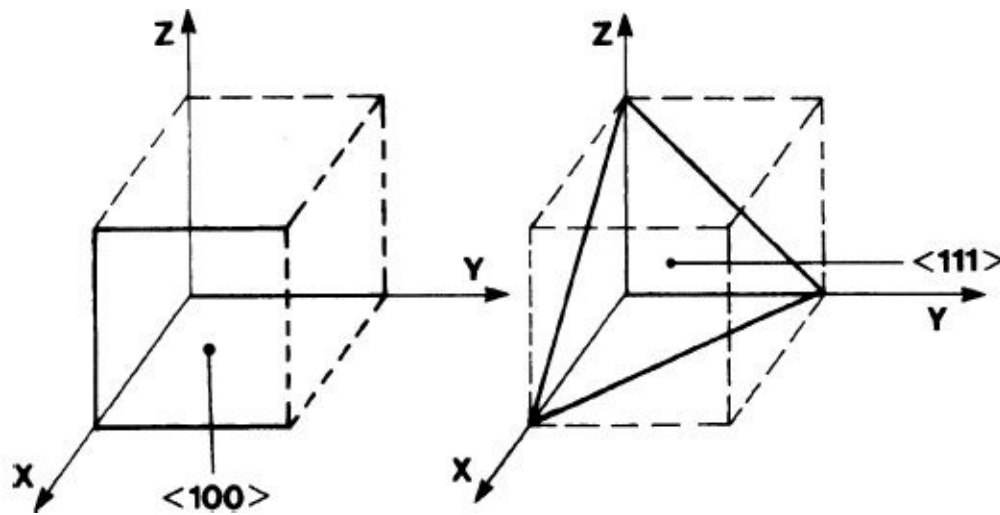


FIGURE 3.7 Crystal planes.

Crystal planes are identified by a series of three numbers known as Miller indices. Two simple cubic unit cells nestled into the origin of an XYZ coordinate system are shown in [Fig. 3.7](#). The two most common orientations used for silicon wafers are the $\langle 100 \rangle$ and the $\langle 111 \rangle$ planes. The plane designations are verbalized as the one-oh-oh plane and the one-one-one plane. The brackets indicate that the three numbers are Miller indices.

The wafers that are oriented $\langle 100 \rangle$ are used for fabricating metal oxide silicon (MOS) devices and circuits, while the wafers that are oriented $\langle 111 \rangle$ are used for bipolar devices and circuits. GaAs wafers are also cut along the $\langle 100 \rangle$ planes of the crystal.

Note that the $\langle 100 \rangle$ plane in [Fig. 3.7](#) has a square shape, while the $\langle 111 \rangle$ plane is triangular in shape. These orientations are revealed when wafers are broken as shown in [Fig. 3.8](#). The $\langle 100 \rangle$ wafers break into quarters or with right angle (90°) breaks. The $\langle 111 \rangle$ wafers break into triangular pieces.

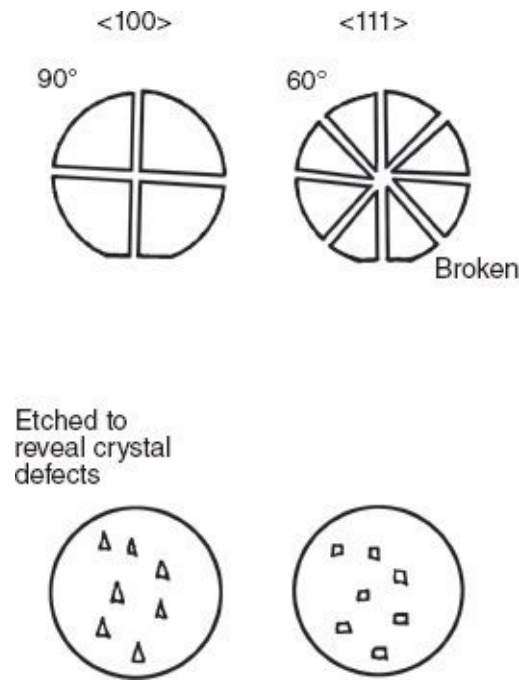


FIGURE 3.8 Wafer orientation indicators.

Crystal Growth

Semiconductor wafers are sliced from large crystals of the semiconducting material. These crystals, also called *ingots*, are grown from chunks of the intrinsic material, which have a polycrystalline structure and are undoped. The process of converting the polycrystalline chunks to a large crystal of single-crystal structure, with the correct orientation and the proper amount of N-or P-type, is called *crystal growing*.

Three different methods are used to grow crystals: the Czochralski (CZ), liquid encapsulated Czochralski, and float-zone techniques.

Czochralski Method

The majority of silicon crystals are grown by the CZ method ([Fig. 3.9](#)). The equipment consists of a quartz (silica) crucible that is heated by surrounding coils that carry radio frequency (RF) waves or by electric heaters. The crucible is loaded with chunks of polycrystalline of the semiconductor material and small amounts of dopant. The dopant material is selected to create either an N-type or P-type crystal. First, the poly and dopants are heated to the liquid state at 1415°C ([Fig. 3.9](#)). Next, a seed crystal is positioned to just touch the surface of the liquid material (called the *melt*). The seed is a small crystal that has the same crystal orientation that is required in the finished crystal ([Fig. 3.10](#)). Seeds can be produced by chemical vapor techniques. In practice, they are pieces of previously grown crystals reused as seeds.

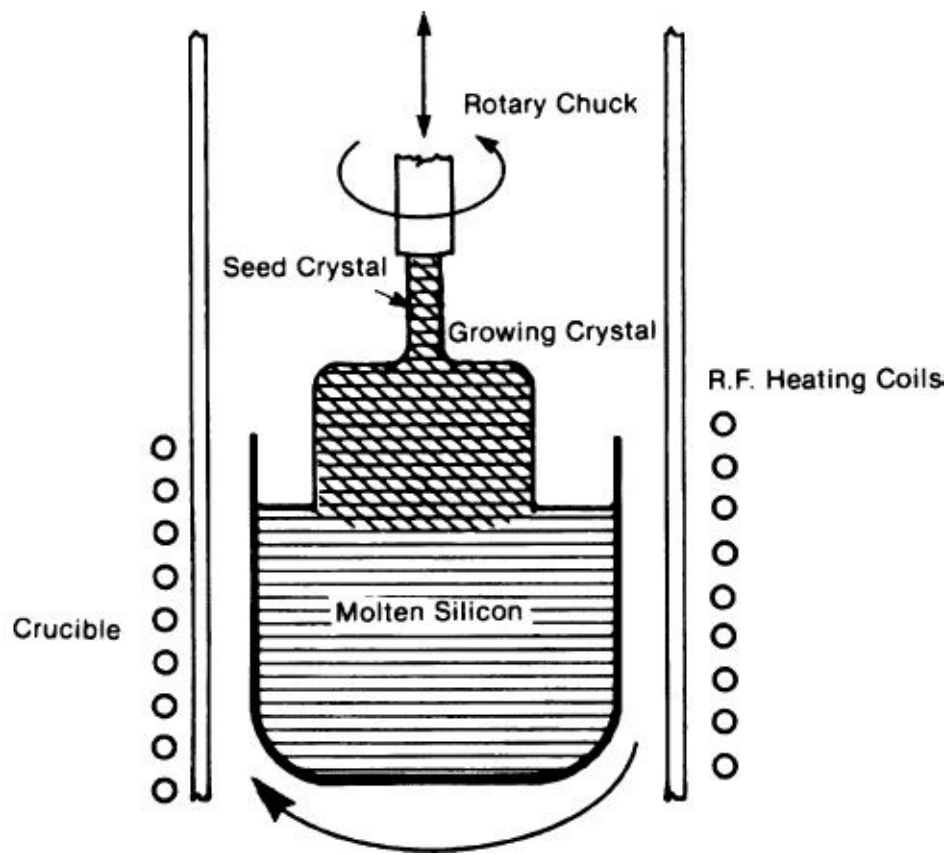


FIGURE 3.9 Czochralski crystal-growing system.

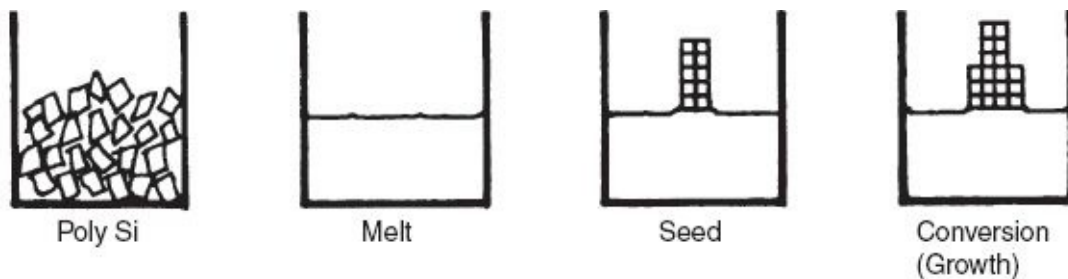


FIGURE 3.10 Crystal growth from a seed.

Crystal growth starts as the seed is slowly raised above the melt. The surface tension between the seed and the melt causes a thin film of the melt to adhere to the seed and then to cool. During the cooling, the first layer atoms from the melted semiconductor material orient themselves to the crystal structure of the seed. Atoms of successive layers continue to replicate the orientation of the seed crystal. The net effect is that the crystal orientation of the seed is propagated in the growing crystal. The dopant atoms in the melt become incorporated into the growing crystal, creating an N-or P-type crystal.

To achieve doping uniformity, crystal perfection, and diameter control, the pull rate is controlled and the seed and crucible are rotated in opposite directions during the entire crystal-growing process. Process control requires a complicated feedback

system integrating the parameters of rotational speed, pull speeds, and the melt temperature.

The crystal is pulled in three sections. First a thin neck is formed, followed by the body of the crystal and ending with a blunt tail. The CZ method is capable of producing crystals several feet in length and with diameters up to 450 mm (18 in). A crystal of 450-mm wafers will weigh some 800 kg and take three days to grow.

Heavier crystals can result in fracturing the small diameter (about 4 mm) seed.⁴ One solution is to start the growth with a process called *dash necking*. In the beginning growth stage, a thickened section is grown. It provides the mechanical strength to support the larger crystal (see [Fig. 3.11](#)).

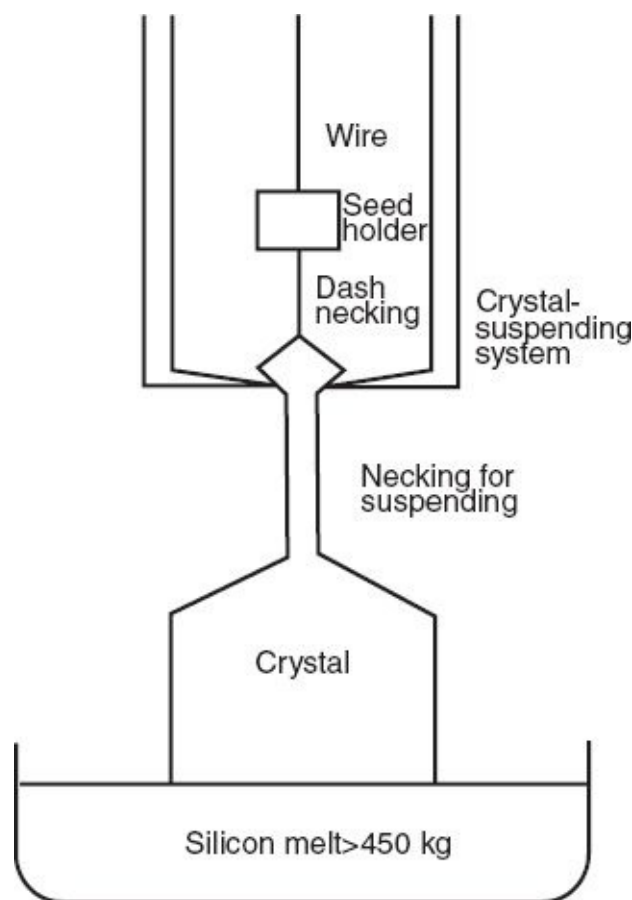


FIGURE 3.11 Dash necking for large-diameter crystals.

Liquid-Encapsulated Czochralski

Liquid-encapsulated czochralski (LEC) crystal growing⁵ is used to grow gallium arsenide crystals. This process is essentially the same as the standard CZ process, but with a major modification. The modification is required because of the evaporative property of the arsenic in the melt. At the crystal growing temperature, the gallium and arsenic react, and the arsenic can evaporate, resulting in a

nonuniform crystal.

Two solutions to the problem are available. One is to pressurize the crystal-growing chamber to suppress the evaporation of the arsenic. The other is the LEC process ([Fig. 3.12](#)). LEC uses a layer of boron tri-oxide (B_2O_3) floating on top of the melt to suppress the arsenic evaporation. In this method, a pressure of about 1 atm is required in the chamber.

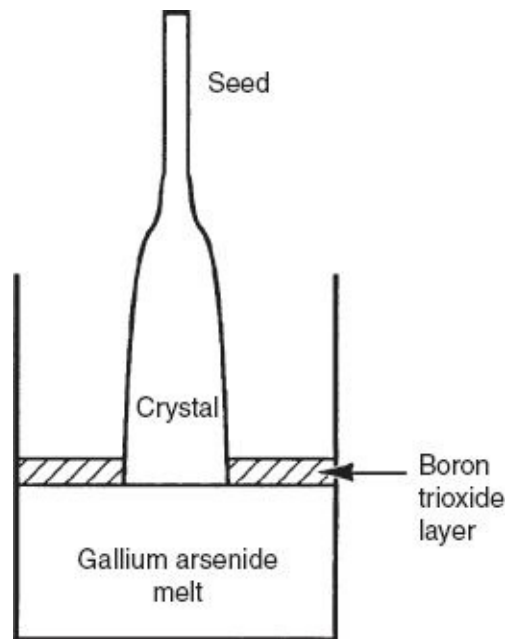


FIGURE 3.12 LEC system of crystal growth.

Float Zone

Float-zone crystal growth is one of the several processes explained in this text that were developed early in the history of the technology and are still used for special needs.⁶ A drawback to the CZ method is the inclusion of oxygen that comes from the crucible. For some devices, higher levels of oxygen are intolerable. For these special cases, the crystal might be grown by the float-zone technique, which produces a lower oxygen content crystal.

Float-zone crystal growth ([Fig. 3.13](#)) requires a bar of the polysilicon and dopants that have been cast in a mold. The seed is fused to one end of the bar and the assemblage is placed in the crystal grower. Conversion of the bar to a single-crystal orientation starts when an RF coil heats the interface region of the bar and seed. The coil is then moved along the axis of the bar, heating it to the liquid point, a small section at a time. Within each molten region, the atoms align to the orientation starting at the seed end. Thus, the entire bar is converted to a single crystal with the orientation of the starting seed.

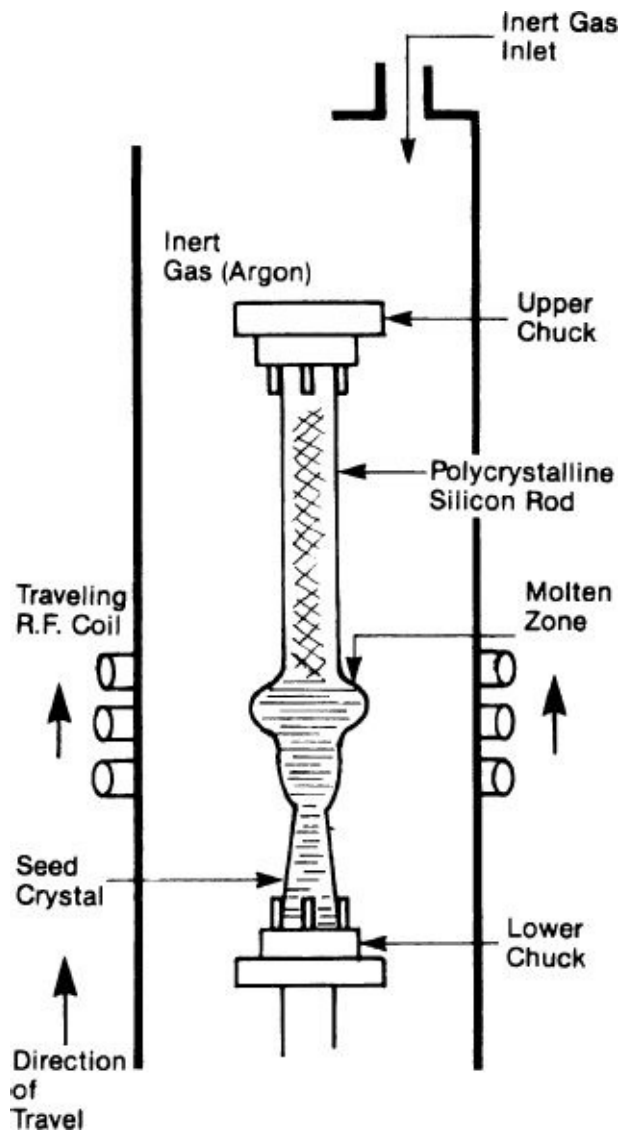


FIGURE 3.13 Float-zone crystal-growing system.

Float-zone crystal growing cannot produce the large diameters that are obtainable with the CZ process, and the crystals have a higher dislocation density. But the absence of a silica (silicon) crucible yields higher-purity crystals with lower oxygen content. Lower oxygen allows crystals with higher purity that find use in semiconductor devices such as power thyristors and rectifiers. The two methods are compared in [Fig. 3.14](#).

PARAMETER	CZ	FLOAT ZONE
Large Crystal	Yes	Difficult
Cost	Lower	
Dislocations	$0 - 10^4/\text{cm}^2$	$10^3 - 10^5/\text{cm}^2$
Resistivity	Up to 100 ohm-cm	2000 ohm-cm Max.
Radial		
Resistivity	5 - 10%	5 - 10%
Oxygen Content	$10^{16} - 10^{18}$ atoms/cm ³	0 - Very Low

FIGURE 3.14 Comparison of CZ and float-zone crystal-growing methods.

Crystal and Wafer Quality

Semiconductor devices require a high degree of crystal perfection. But even with the most sophisticated techniques, a perfect crystal is unobtainable. The imperfections, called *crystal defects*, result in process problems by causing uneven silicon dioxide film growth, poor epitaxial film deposition, uneven doping layers in the wafer, and other problems. In finished devices, the crystal defects cause unwanted current leakage and may prevent the devices from operating at required voltages. There are four major categories of crystal defects:

1. Point defects
2. Dislocations
3. Growth defects
4. Impurities

Point Defects

Point defects come in two varieties. One is when contaminants in the crystal become jammed in the crystal structure, causing strain. The second is known as a vacancy. In this situation, there is an atom missing from a location in the structure ([Fig. 3.15](#)).

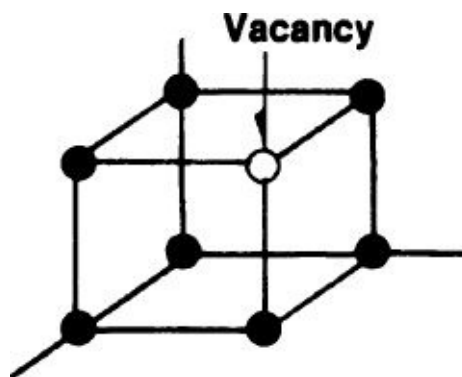


FIGURE 3.15 Vacancy crystal defect.

Vacancies are natural phenomena that occur in every crystal. Unfortunately, vacancies occur whenever a crystal or wafer is heated and cooled, such as in the fabrication process. The minimization of vacancies is one of the driving forces behind the desire for low-temperature processing.

Dislocations

Dislocations are misplacements of the unit cells in a single crystal. They can be imagined as an orderly pile of sugar cubes with one of the cubes slightly out of alignment with the others. Dislocations occur from growth conditions and lattice strain in the crystal. They also occur in wafers from physical abuse during the fab process. A chip or abrasion of the wafer edge serves as a lattice strain site that can generate a line of dislocations, which progresses into the wafer interior with each subsequent high-temperature processing of the wafer. Wafer dislocations are revealed by a special etch of the surface. A typical wafer has a density of less than 500 dislocations per square centimeter.

Etched dislocations appear on the surface of the wafer in shapes indicative of their crystal orientation. $\langle 111 \rangle$ wafers etch into triangular dislocations, and $\langle 100 \rangle$ wafers show “squarish” etch pits ([Fig. 3.8](#)).

Growth Defects

During crystal growth, certain conditions can result in structural defects. One is *slip*, which refers to the slippage of the crystal along the crystal planes ([Fig. 3.16](#)). Another problem is *twinning*. This is a situation in which the crystal grows in two different directions from the same interface. Both of these defects are cause for rejection of the crystal.

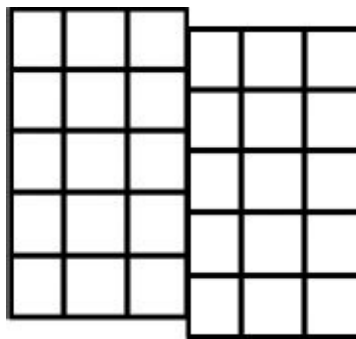


FIGURE 3.16 Crystal slip.

Impurities

Besides unwanted impurities from materials and handling, the CZ process itself

adds two crystal impurities. One is oxygen from the silica crucible. The other is carbon from the graphite in the hot zone. Oxygen is electrically active in the resultant wafers and circuits. The carbon can enhance oxygen precipitation.²

Wafer Preparation

End Cropping

After removal from the crystal grower, the crystal goes through a series of steps that result in the finished wafer. First is the cropping off of the crystal ends with a saw.

Diameter Grinding

During crystal growth, there is a diameter variation over the length of the crystal ([Fig. 3.17](#)). Wafer-fabrication processing, with its variety of wafer holders and automatic equipment, requires tight diameter control to minimize warped and broken wafers from the handling tools.

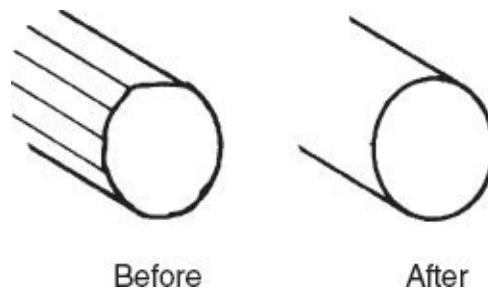


FIGURE 3.17 Crystal diameter grinding.

Diameter grinding is a mechanical operation performed in a *centerless grinder*. This machine grinds the crystal to the correct diameter without the necessity of clamping it into a lathe-type grinder with a fixed center point—although lathe-type grinders are used.

Crystal Orientation, Conductivity, and Resistivity Check

Before the crystal is submitted to the wafer preparation steps, it is necessary to determine whether it meets orientation and resistivity specifications.

The crystal orientation ([Fig. 3.18](#)) is determined by either X-ray diffraction or

collimated light refraction. In both methods, an end of the crystal is etched or polished to remove saw damage. Next, the crystal is mounted in the refraction apparatus and the X-rays or collimated light reflected off the crystal surface onto a photographic plate (X-rays) or screen (collimated light). The pattern formed on the plate or screen is indicative of the crystal plane (orientation) of the grown crystal. The pattern shown in [Fig. 3.18](#) is representative of a $\langle 100 \rangle$ orientation.

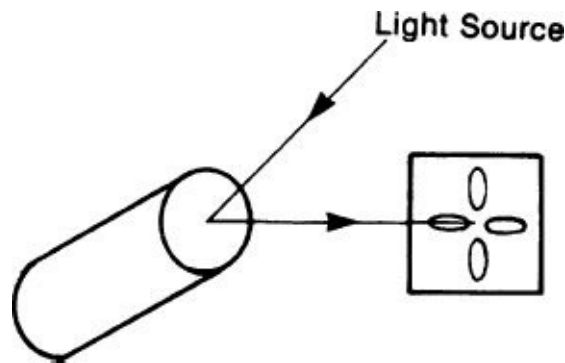


FIGURE 3.18 Crystal orientation determination.

Most crystals are purposely grown several degrees off the major $\langle 111 \rangle$ or $\langle 100 \rangle$ plane. This off-orientation provides several benefits in wafer-fabrication processing, particularly ion implantation. The reasons are covered in the applicable process chapters.

The crystal is positioned on a slicing block to ensure that the wafers will be cut from the crystal in the correct orientation.

Because each crystal is doped, an important electrical check is the conductivity type (N or P) to ensure that the right dopant type was used. A hot-point probe connected to a polarity meter is used to generate holes or electrons (depending on the type) in the crystal. The conductivity type is displayed on the meter.

The amount of dopant put into the crystal is determined by a resistivity measurement using a four-point probe. See [Chap. 13](#) for a description of this measurement technique. The curves presented in [Chap. 2](#) ([Fig. 2.7](#)) show the relationship between resistivity and doping concentration for N- and P-type silicons.

The resistivity is checked along the axis of the crystal due to dopant variation during the growing process. This variation results in wafers that fall into several resistivity specification ranges. Later in the process, the wafers can be grouped by the resistivity range to meet customer specifications.

Grinding Orientation Indicators

Once the crystal is oriented on the cutting block, a flat or notch is ground along the axis ([Fig. 3.19](#)). The first flat is called the *major flat*. The position of the flat or

notch is along one of the major crystal planes, as determined by the orientation check.

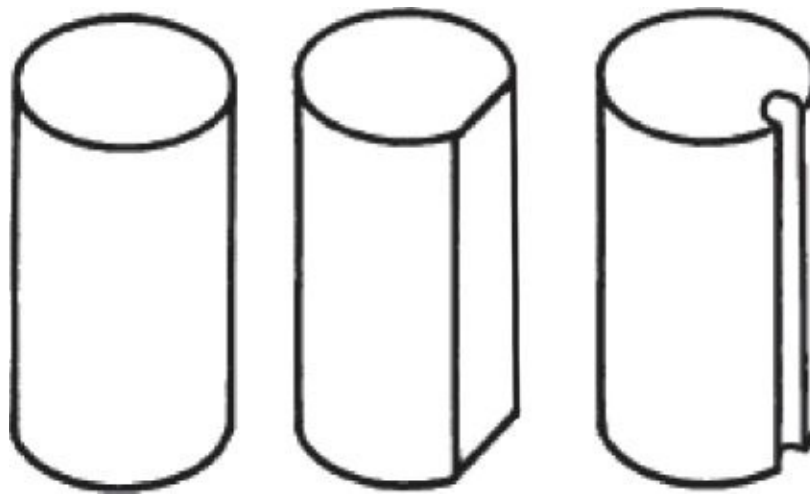


FIGURE 3.19 Crystal flat and notch grinding.

During the fabrication process, the flat functions as a visual reference to the orientation of the wafer. It is used to place the first pattern mask on the wafer so that the orientation of the chips is always to a major crystal plane.

In most crystals, there is a second, smaller, secondary flat ground on the edge. The position of the secondary flat to the major flat is a code that indicates both the orientation and conductivity type of the wafer. The code is shown in [Fig. 3.20](#).

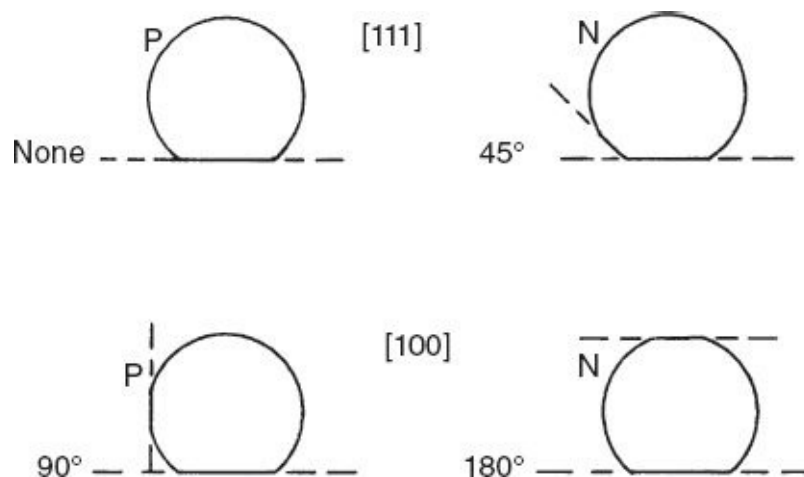
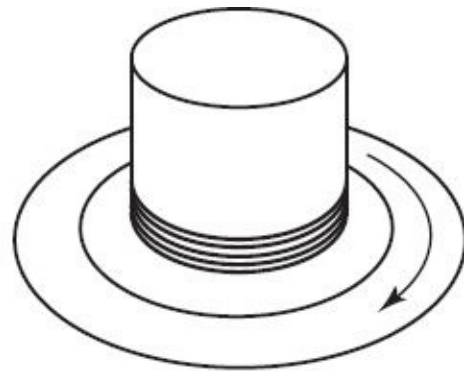


FIGURE 3.20 Wafer flat locations.

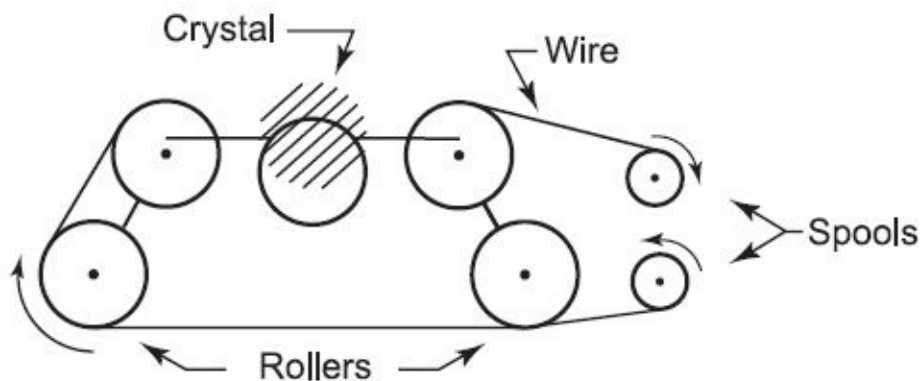
For larger wafer diameters, a notch is ground on the crystal to indicate the wafer crystal orientation ([Fig. 3.19](#)).

Wafer Slicing

The wafers are sliced from the crystal with the use of diamond-coated inside diameter (ID) saws or wire saws ([Fig. 3.21](#)). ID saws are thin circular sheets of steel with a hole cut out of the center. The inside of the hole is the cutting edge and is coated with diamonds. An inside diameter saw has rigidity, but without being very thick. These factors reduce the kerf (cutting width) size, which in turn prevents sizable amounts of the crystal being wasted by the slicing process.



a. Inside Diameter Diamond Saw



b. Wire Saw

FIGURE 3.21 Wafer slicing.

For larger-diameter wafers (> 300 mm), wire saws are used to ensure flat surfaces with little tapering and with a minimal amount of “kerf” loss.

Wafer Marking

Large-area wafers represent a high value in the wafer-fabrication process.

Identifying them is necessary to prevent misprocessing and to maintain accurate tracking. Bar codes or a data-dot matrix code, generally in a square pattern, are laser inscribed on the wafer⁸ (Fig. 3.22). Laser dots are the agreed upon method for 300-mm and larger wafers.

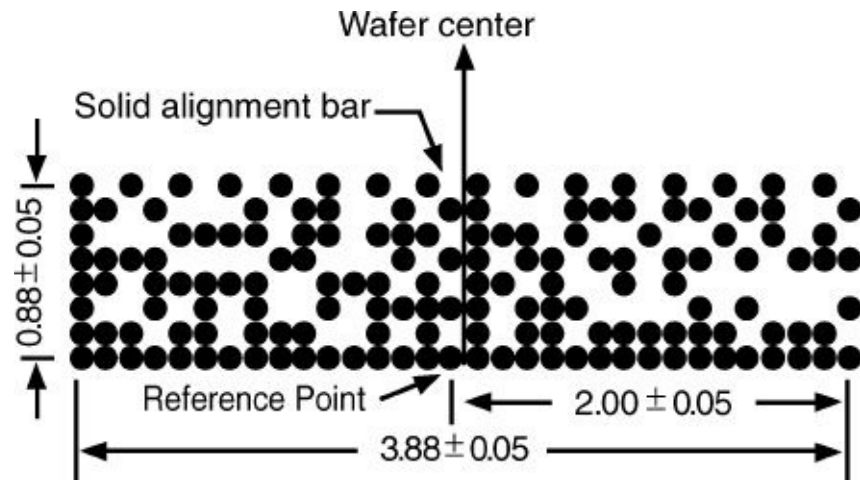


FIGURE 3.22 Laser dot coding. (Reprinted from the Jan. 1998 edition of Solid State Technology, Copyright 1998 by PennWell Publishing Company.)

Rough Polish

The surface of a semiconductor wafer has to be free of irregularities and saw damage and must be absolutely flat. The first requirement comes from the very small dimensions of the surface and subsurface layers making up the device. Newer devices have dimensions (feature sizes) in the nanometer (nm) range. To get an idea of the relative dimensions of a semiconductor device, imagine the cross-section in Fig. 3.23 (as tall as a house wall), about 8 ft (2.4 m). On that scale, the working layers on the top of the wafer would all exist within the top 1 or 2 in (25 or 50 mm) or less.

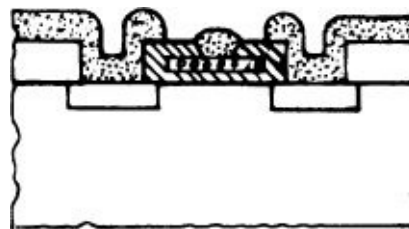


FIGURE 3.23 Cross-section of an MOS transistor.

Flatness is an absolute requirement for small-dimension patterning (see Chap. 11). The advanced patterning processes project the required pattern image onto the wafer surface. If the surface is not flat, the projected image will be distorted just as a

film image will be out of focus on a non-flat screen.

The flattening and polishing process proceeds in two steps: rough polish and chemical/mechanical polishing ([Fig. 3.24](#)). Rough polishing is a conventional abrasive slurry lapping process, but it is fine-tuned to semiconductor requirements. A primary purpose of the rough polish is to remove the surface damage left over from the wafer-slicing process.

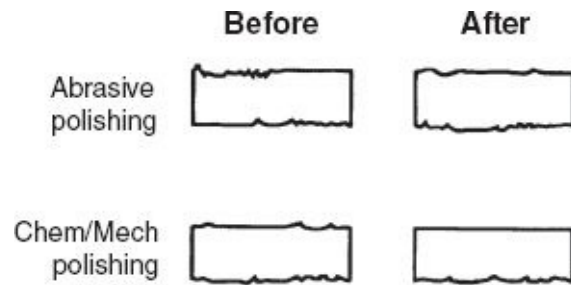


FIGURE 3.24 Abrasive and chemical-mechanical surface polishing.

Chemical Mechanical Polishing

The final polishing step is a combination of chemical etching and mechanical buffing called chemical mechanical polishing (CMP). The wafers are mounted on rotating holders and lowered onto a pad surface rotating in the opposite direction. The pad material is generally a cast and sliced polyurethane with a filler or a urethane-coated felt. A slurry of silica (glass) suspended in a mild etchant, such as potassium or ammonium hydroxide, is fed onto the polishing pad.

The alkaline slurry chemically grows a thin layer of silicon dioxide on the wafer surface. The buffing action of the pad mechanically removes the oxide in a continuous action. High points on the wafer surface are removed until an extremely flat surface is achieved. If a typical semiconductor wafer surface was extended to 10,000 ft (the length of a typical airport runway), it would vary about plus or minus 2 in or less over its entire length.

Achieving the extreme flatness requires the specification and control of the polishing time, the pressure on the wafer and pad, the speed of rotation, the slurry particle size, the slurry feed rate, the chemistry (pH) of the slurry, and the pad material and conditions.

Chemical mechanical polishing is one of the techniques developed by the industry that has allowed production of larger wafers. CMP is used in the wafer-fabrication process to flatten wafer surfaces after the buildup of new layers creates uneven surfaces. In this application, the CMP process is used for planarization of the surface. A detailed explanation of this use of CMP appears in [Chap. 10](#).

Backside Processing

In most cases, only the front side of the wafer goes through the extensive CMP. The backs may be left rough or etched to a bright appearance. For some device use, the backs may receive a special process to induce crystal damage, called *backside damage*. Backside damage causes the growth of dislocations that radiate up into the wafer. These dislocations can act as a trap for mobile ionic contamination introduced into the wafer during the fab process. The trapping phenomenon is called *gettering* ([Fig. 3.25](#)). The backside processes include sandblasting or a deposition of a polysilicon layer or silicon nitride layer on the back.

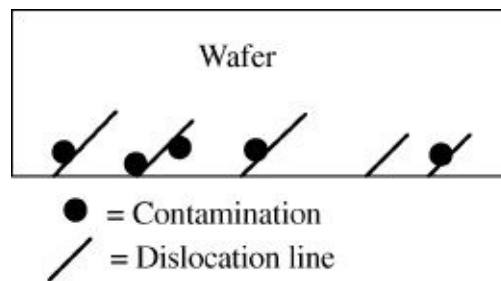


FIGURE 3.25 Trapping.

Double-Sided Polishing

One of the many demands on larger-diameter wafers is flat and parallel surfaces. Most manufacturers of 300-mm diameter wafers employ double-sided polishing to achieve flatness specifications of 0.25 to 0.18 μm over 25 ? 25 mm sites.⁹ A downside is that all further processing must employ handling techniques that do not scratch or contaminate the backside.

Edge Grinding and Polishing

Edge grinding is a mechanical process that leaves the wafer with a rounded edge ([Fig. 3.26](#)). Chemical polishing is employed to further create an edge that minimizes edge chipping and damage during fabrication that can result in wafer breakage or serve as the nucleus for dislocation lines that propagate into the chips near the edge of the wafer.



FIGURE 3.26 Wafer-edge grinding.

Wafer Evaluation

Before packing, the wafers (or samples) are checked for a number of parameters as specified by the customer. [Figure 3.27](#) illustrates a typical wafer specification.

Diameter	300 mm, +/- .02 mm
Thickness	775 um +/- 25 um
Orientation	{ 100 } =/- 2°
Resistivity	Greater than 1 ohm-em
Oxygen	20–31 ppma
Carbon	Less than 0.2 ppma

FIGURE 3.27 Typical 300-mm wafer specification.

Primary concerns are surface problems such as particulates, stain, and haze. These problems are detected with the use of high-intensity lights or automated inspection machines.

Oxidation

Silicon wafers may be oxidized before shipment to the customer. The silicon dioxide layer serves to protect the wafer surface from scratches and contamination during shipping. Most companies start the wafer-fabrication process with an oxidation step, and buying the wafers with an oxide layer saves a manufacturing step. Oxidation processes are explained in [Chap. 7](#).

Packaging

While much effort goes into producing a high-quality and clean wafer, the quality can be lost during shipment to the customer or, worse, from the packaging method itself. Therefore, there is a very stringent requirement for clean and protective packaging. The packaging materials are of nonstatic, nonparticle-generating materials, and the equipment and operators are grounded to drain off static charges that attract small particles to the wafers. Wafer packaging takes place in cleanrooms.

Wafer Types and Uses

These processes are geared to produce *prime wafers*, which are the host for producing chips and circuits in the wafer-fabrication process. In addition there is a need for various types of *test or monitor wafers*. These are used in place of

expensive prime wafers to monitor and evaluate the outcomes of the process steps. These are *mechanical test wafers* and *process test wafers*.¹⁰

Mechanical test wafers are used to test and verify operational aspects of equipment and handling systems. Process test wafers (also known as monitor wafers) go into the process steps with the prime wafers and/or through the process modules. Prime wafers cannot be used to test and control the outcome of a single-process step. For example, measuring the thickness of the 15 layer in the process with 14 layers already on the wafer is impossible. Hence blank process test wafers are needed.

Reclaim Wafers

Wafers are rejected out of the process for many reasons, usually for not meeting process specifications. Assuming they are not physically damaged, they may be reclaimed for use as test wafers. A combination of chemical and CMP is used. Top-added layers and the top layer of the wafer are removed. This creates a new wafer surface suitable for test wafer use. After layer removal, a reclaimed wafer will go through the same wafer-cleaning processes as a prime wafer.

Engineered Wafers (Substrates)

Increasingly, wafer-fabrication companies are asking for wafer manufacturers to supply wafers with deposited top-side layers, such as epitaxial silicon. Other wafer products include silicon deposited on insulators (SOI and SOS) such as sapphire or diamond ([Chap. 12](#)).

Review Topics

Upon completion of this chapter, you should be able to: 1. Explain the difference between crystalline and noncrystalline materials.

2. Explain the differences between a polycrystalline and a single crystalline material.

3. Draw a diagram of the two major wafer crystal orientations used in semiconductor processing.

4. Explain the Czochralski, float zone, and liquid crystal encapsulated Czochralski methods of crystal growing.

5. Draw a flow diagram of the wafer preparation process.

6. Explain the use and meaning of the flats or notches ground on wafers.

7. Describe the benefits in the wafer-fab process that come from edge-

rounded wafers.

8. Describe the benefits in the wafer-fab process that come from flat and damage-free wafer surfaces.

References

1. Arensman, R., "One-Stop Automation," *Electronic Business*, Jul. 2002:54.
2. Watanabe, M., and Kramer, S., "450-mm Silicon: An Opportunity and Wafer Scaling," *Electro Chemical Interface*, Winter 2006:28 ff.
3. Wolf, S., *Microchip Manufacturing*, 2004, Lattice Press, Sunset Beach, CA:148.
4. Lin, W., and Huff, H., *Handbook of Semiconductor Manufacturing Technology*, 2nd ed., CRC Press, Hoboken, NJ:3–42.
5. Williams, R. E., *Gallium Arsenide Processing Techniques*, 1984, Artech House Inc., Dedham, MA:37.
6. Silicon Consultant, *Single Crystal Growth by Float Zone*, www.siliconconsultant.com/SIcrtgr.htm, (May, 2013).
7. Lin, W., and Huff, H., *Handbook of Semiconductor Manufacturing Technology*, 2nd ed., CRC Press, Hoboken, NJ:3–9.
8. Brunkhorst, S. J., and Sloat, D. W., "The Impact of the 300-mm Transition on Silicon Wafer Suppliers," *Solid State Technology*, Jan. 1998:87.
9. Ibid.
10. Advantiv Product List, *Bare Silicon Wafers*, www.advantivtech.com/wafers/silicon (May, 2013).

CHAPTER 4

Overview of Wafer Fabrication and Packaging

Introduction

This chapter will introduce the four basic processes used in the wafer fabrication to form the electrical elements of an integrated circuit (IC) in and on the wafer surface. It includes the starting activity of circuit design through to the production of photomasks and reticles. Wafer and chip features and terminology are detailed. A flow diagram shows the steps to build a simple semiconductor device.

At the end of the wafer-fabrication process, functioning chips are advanced to the packaging stage. There are many options and processes from mounting a bare chip directly onto a circuit board to the stacking (3-D) of multiple chips (die) in the same package. The basic steps and package options are presented.

Goal of Wafer Fabrication

The stages of microchip manufacturing are materials preparation, crystal growth or wafer preparation, wafer fabrication and electrical sort, and packaging or final test. The first two stages have been explored in [Chap. 3](#). In this chapter, the fundamentals of Stage 3, wafer fabrication, are introduced. [Chapters 5](#) to [14](#) cover the specific fabrication processes and technologies. The details of electrical wafer sort and packaging are explained in [Chap. 18](#).

Wafer fabrication is the manufacturing process used to create semiconductor devices and circuits in and on a wafer surface. The polished starting wafers come into fabrication with blank surfaces and exit with the surface covered with hundreds of completed chips ([Fig. 4.1](#)).

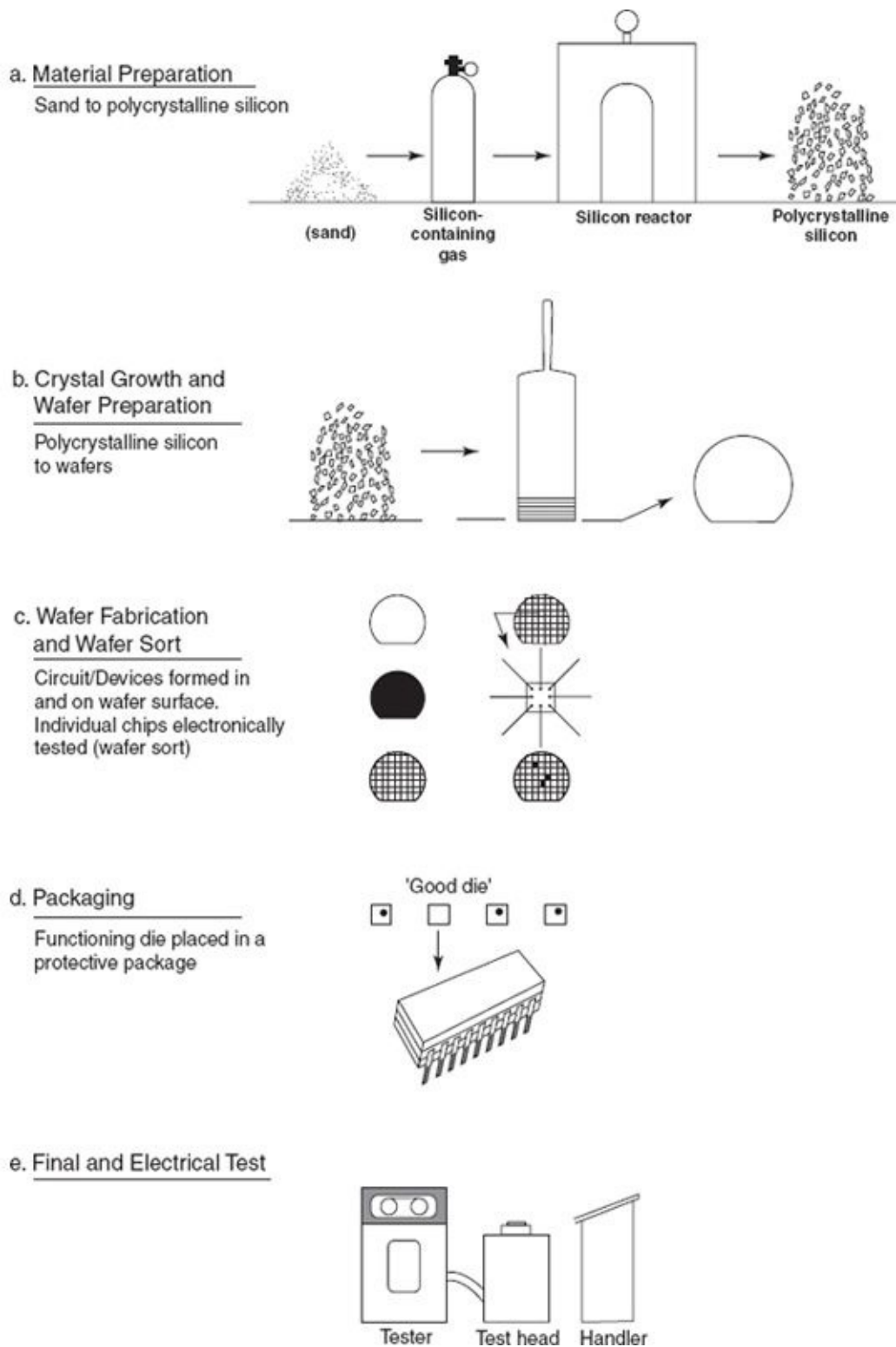


FIGURE 4.1 Stages of semiconductor production.

Wafer Terminology

A completed wafer is illustrated in [Fig. 4.2](#). The regions of a wafer surface are:

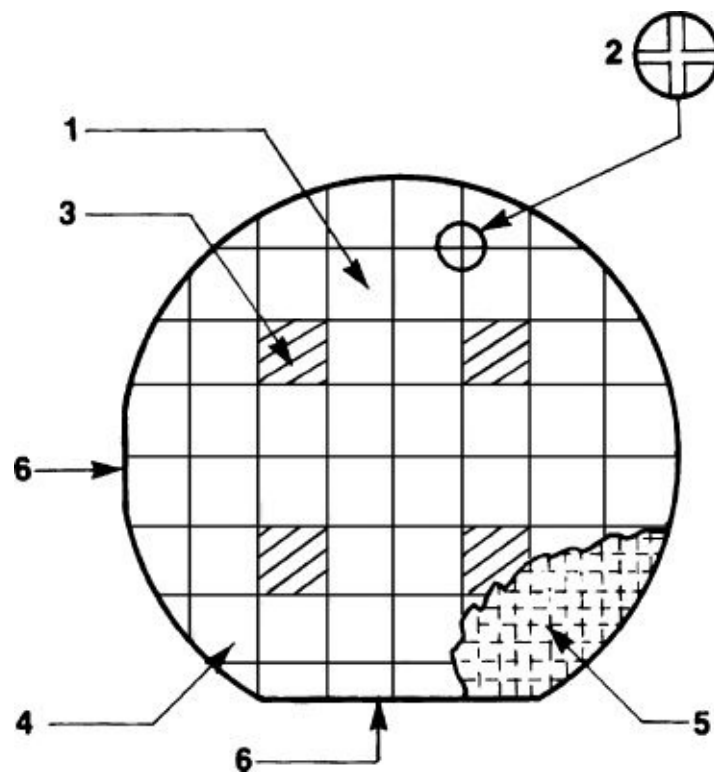


FIGURE 4.2 Wafer terminology.

1. Chip, die, device, circuit, microchip, or bar. All of these terms are used to identify the microchip patterns covering the majority of the wafer surface.

2. Scribe lines, saw lines, streets, and avenues. These areas are spaces between the chips that allow separation of the chip from the wafer. Generally, the scribe lines are blank, but some companies place alignment targets or electrical test structures (see photomasking) in them.

3. Engineering die and test die. These chips are different from the regular device or circuit die. They contain special devices and circuit elements that allow electrical testing during the fabrication processing for process and quality control.

4. Edge chips. The edges of the wafer contain partial chip patterns that are wasted space. Larger chips on the same diameter wafer result in large numbers of partial chips. Larger-diameter wafers minimize the amount of wasted space from larger chips.

5. Wafer crystal planes. The cutaway section illustrates the crystal structure of the wafer under the circuit layers. The diagram shows that the chip edges are oriented to the wafer crystal structure.

6. Wafer flats/notches. Wafers come from the wafer preparation stage with flats or notches, which indicate the crystal orientation and doping polarity of the wafer. Longer flats are called *major flats*, and shorter flats are called *minor flats*. The depicted wafer has a major and minor flat, indicating that it is a P-

type $\langle 100 \rangle$ oriented wafer (see [Chap. 3](#) for the flat code). The wafers of 300 and 450 mm diameter use notches as crystal orientation indicators. The flats and notches also assist alignment of the wafer in a number of the wafer-fabrication processes.

Chip Terminology

[Figure 4.3](#) depicts a photomicrograph of a Medium Scale Integration (MSI)/bipolar integrated circuit. The level of integration was chosen so that some surface details could be seen. The components of higher-density circuits are so small that they cannot be distinguished on a photomicrograph of the entire chip.

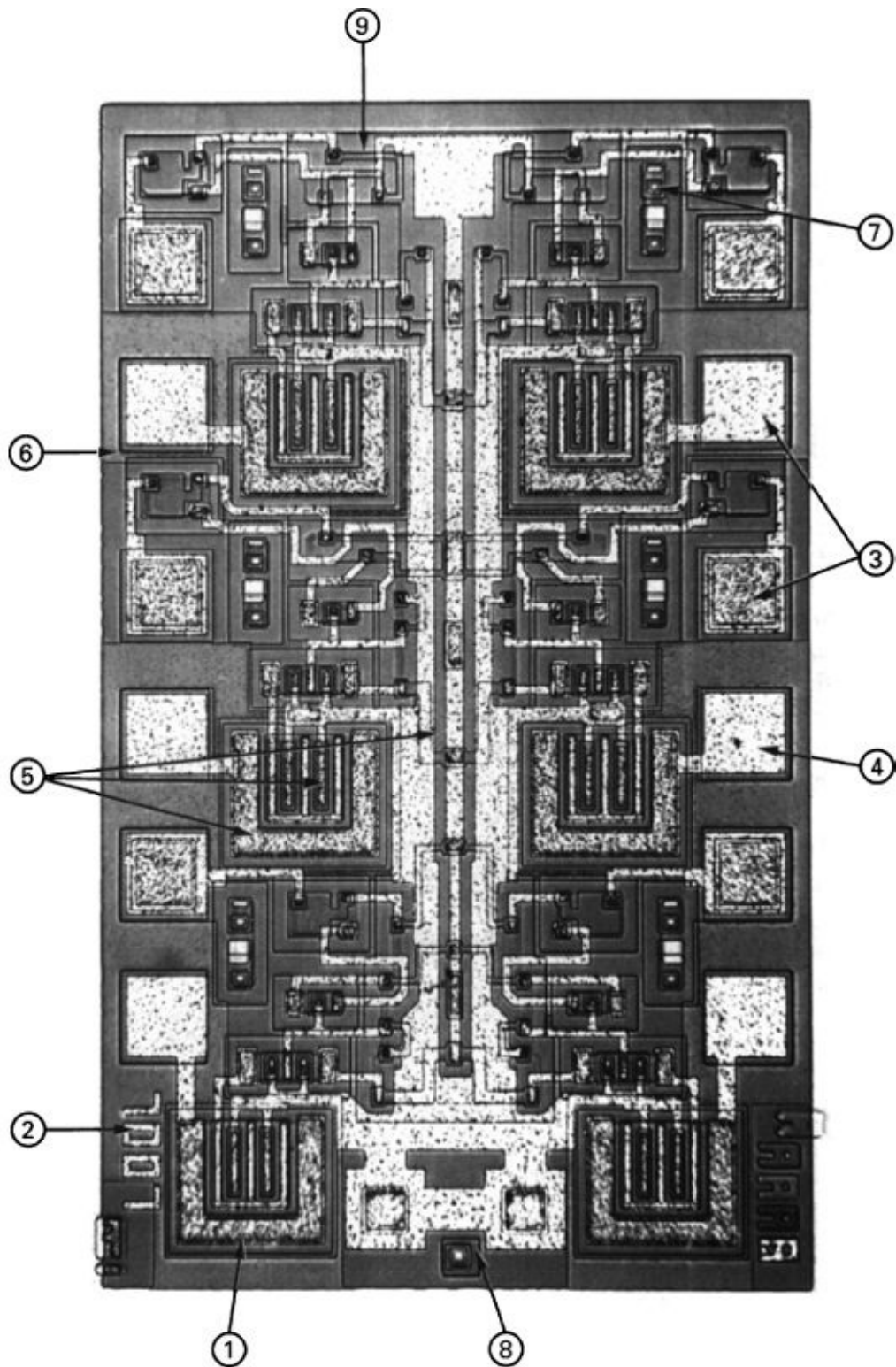


FIGURE 4.3 Chip terminology.

The chip features are:

1. A bipolar transistor
2. The circuit designation number
3. Bonding pads for connecting the chip into a package
4. A piece of contamination on a bonding

pad 5. Metal surface “wiring” lines 6. Scribe (separation) lines 7. Unconnected component 8. Mask-alignment marks 9. Resistor

Basic Wafer-Fabrication Operations

There are many hybrid and integrated circuit designs. ICs are based on a small number of transistor structures (primarily bipolar or metal oxide silicon [MOS] structures, see [Chap. 16](#)) and manufacturing processes. An analogy is the auto industry. This industry produces a wide variety of products, from sedans to bulldozers. Yet the processes of metal forming, welding, painting, and so on, are common to all plants. Within the plant, these basic processes are applied in different ways to produce different products.

The same is true for microchip fabrication. The four basic operations are sequenced to produce specific microchips. The operations are layering, patterning, doping, and heat treatments. [Figure 4.4](#) is a cross-section of a (MOS) silicon gate transistor. It illustrates how these basic operations are used and sequenced to create a real-life semiconductor device.

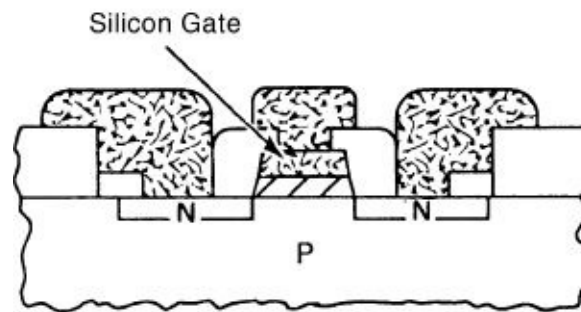


FIGURE 4.4 Basic MOS silicon gate transistor.

Layering

Layering is the operation used to add thin layers to the wafer surface. An examination of the simple MOS transistor structure in [Fig. 4.5](#) shows a number of layers that have been added to the wafer surface. These layers could be insulators, semiconductors, or conductors. They are of different materials and are grown or deposited by a variety of processes.

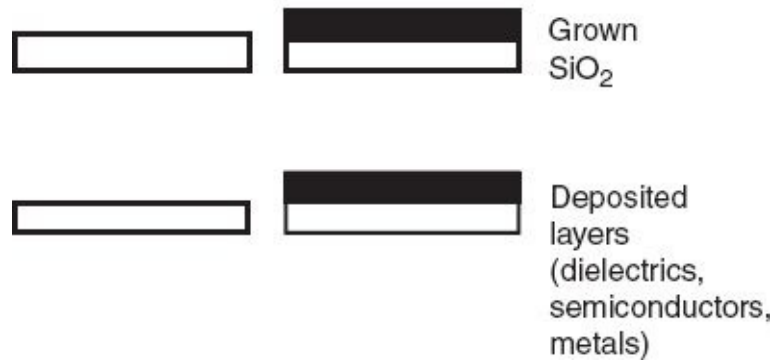


FIGURE 4.5 Layering operations.

Various techniques are used for growing silicon dioxide layers and deposition (Fig. 4.6) of a variety of materials. Common deposition techniques are physical vapor deposition (PVD), chemical vapor deposition (CVD), evaporation, sputtering, molecular beam, epitaxy, molecular beam epitaxy, and atomic layer deposition (ALD). Electroplating is used to deposit gold metallization on high-density integrated circuits. Figure 4.7 lists common layer materials and layering processes. The details of each are explained in the process chapters. The role of the different layers in the structures is explained in Chapter 16.

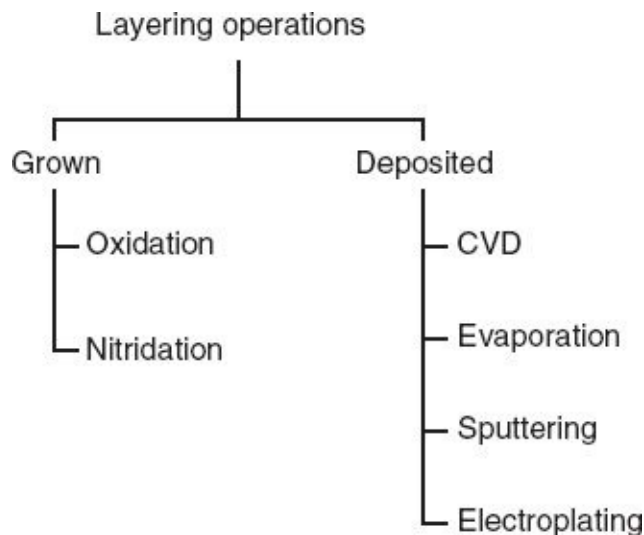


FIGURE 4.6 Layering operations.

Layers	Thermal oxidation	Chemical vapor deposition	Evaporation	Sputtering	Electroplating
Insulators	silicon dioxide	silicon dioxide		silicon dioxide	
		silicon nitride		silicon monoxide	
Semiconductors		epitaxial silicon			
		polycrystalline silicon			
Conductors			aluminum	aluminum	gold
			aluminum alloys	aluminum alloys	copper
			nichrome	tungsten	
			gold	titanium	
				molybdenum	

FIGURE 4.7 Table of layers, processes, and materials.

Patterning

Patterning is the series of steps that results in the removal of selected portions of the added surface layers (Fig. 4.8). After removal, a *pattern* of the layer is left on the wafer surface. The material removed may be in the form of a hole in the layer or just a remaining island of the material.



FIGURE 4.8 Patterning.

The patterning process is known by the names photomasking, masking, photolithography, and microlithography. During the wafer-fabrication process, the various physical *parts* of the transistors, diodes, capacitors, resistors, and metal conduction system are formed in and on the wafer surface. These parts are created one layer at a time by the combination of putting a layer on the surface and using a patterning process to leave a specific shape. The goal of the patterning operation is to create the desired shapes in the exact dimensions (feature size) required by the circuit design, and to locate them in their proper location on the wafer surface and in relation to the other *layers*.

Patterning is the most critical of the four basic operations. This operation sets the critical dimensions of the devices. Errors in the patterning process can cause distorted or misplaced patterns that cause electrical malfunctioning of the device or circuit. Misplacement of the pattern can have the same bad results. Another problem is defects. Patterning is a high-tech version of photography but is performed at incredibly small dimensions. Contamination in the process steps can introduce

defects. This contamination problem is magnified by the fact that patterning operations are performed on the wafer 30 or more times in the course of a state-of-the-art wafer-fabrication process.

Circuit Design

Circuit design is the first step in the creation of a microchip. A circuit designer starts with a block functional diagram of the circuit such as the logic diagram in [Fig. 4.9](#). This diagram lays out the primary functions and operation required of the circuit. Next, the designer translates the functional diagram to a schematic diagram ([Fig. 4.10](#)), showing the number and connection of the various circuit components. Each component is represented by a symbol. Accompanying the schematic diagram are the electrical parameters (voltage, current, resistance, and so forth) required to make the circuit work.

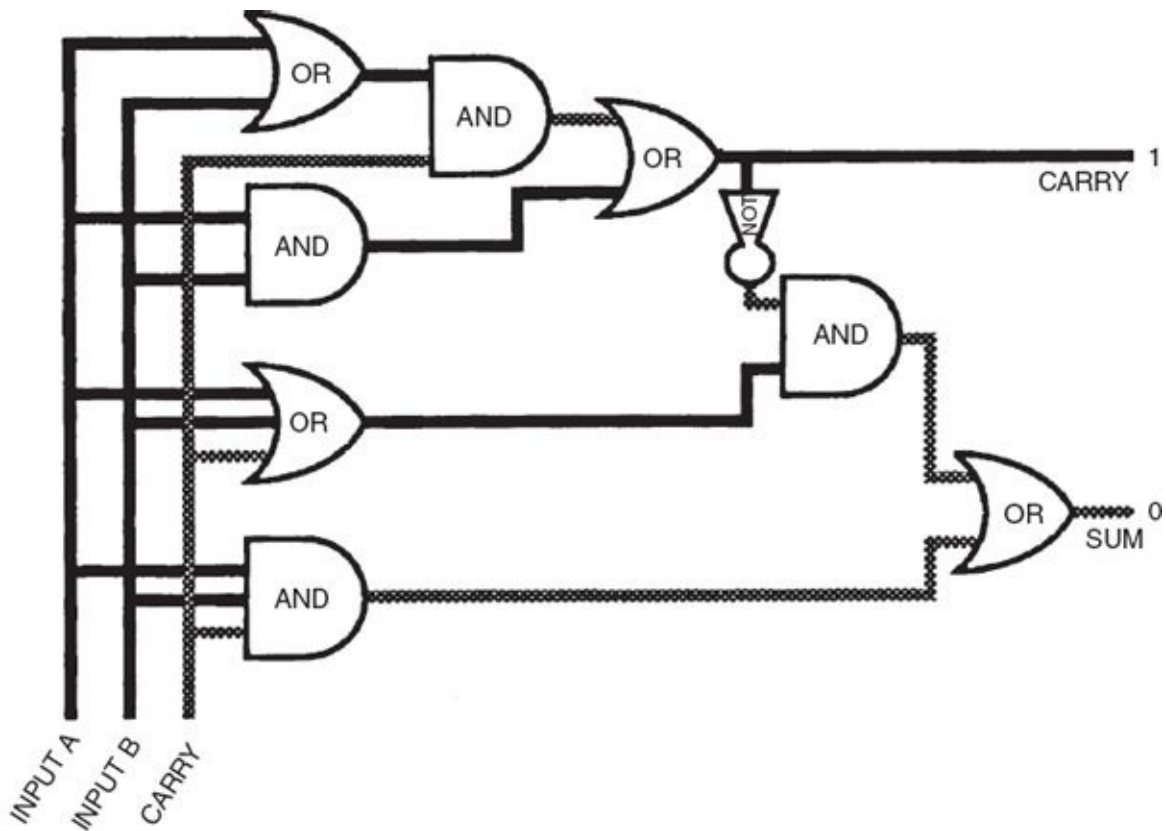


FIGURE 4.9 Example of a functional logic design of a simple circuit.

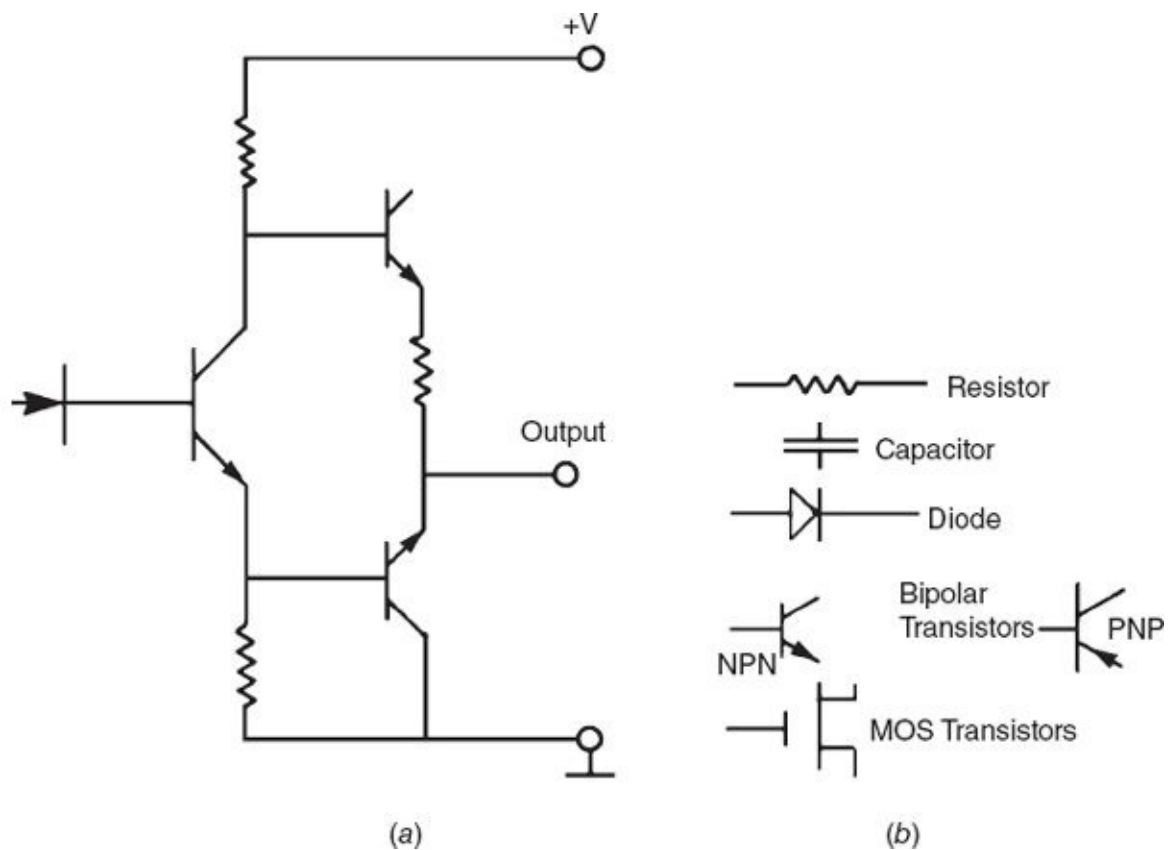


FIGURE 4.10 Example of a circuit schematic diagram with component symbols.

The third step, circuit layout, is unique to semiconductor circuits. Circuit operation is dependent on a number of factors, including material resistivity, material physics, and the physical dimensions of the individual component “parts.” The placement of the parts relative to each other is another factor. All these considerations dictate the physical layout and dimensions of the part, device, or circuit. Layout starts with using sophisticated computer-aided design (CAD) systems to translate each circuit component into its physical shape and size. Through the CAD system, the circuit is built, exactly duplicating the final design. The result is a composite picture of the circuit surface showing all of the sublayer patterns. This drawing is called a *composite drawing* (Fig. 4.11). The composite drawing is analogous to the blueprint of a multistory office building as viewed from the top and showing all of the floors. However, the composite is many times larger than the dimensions of the final circuit.

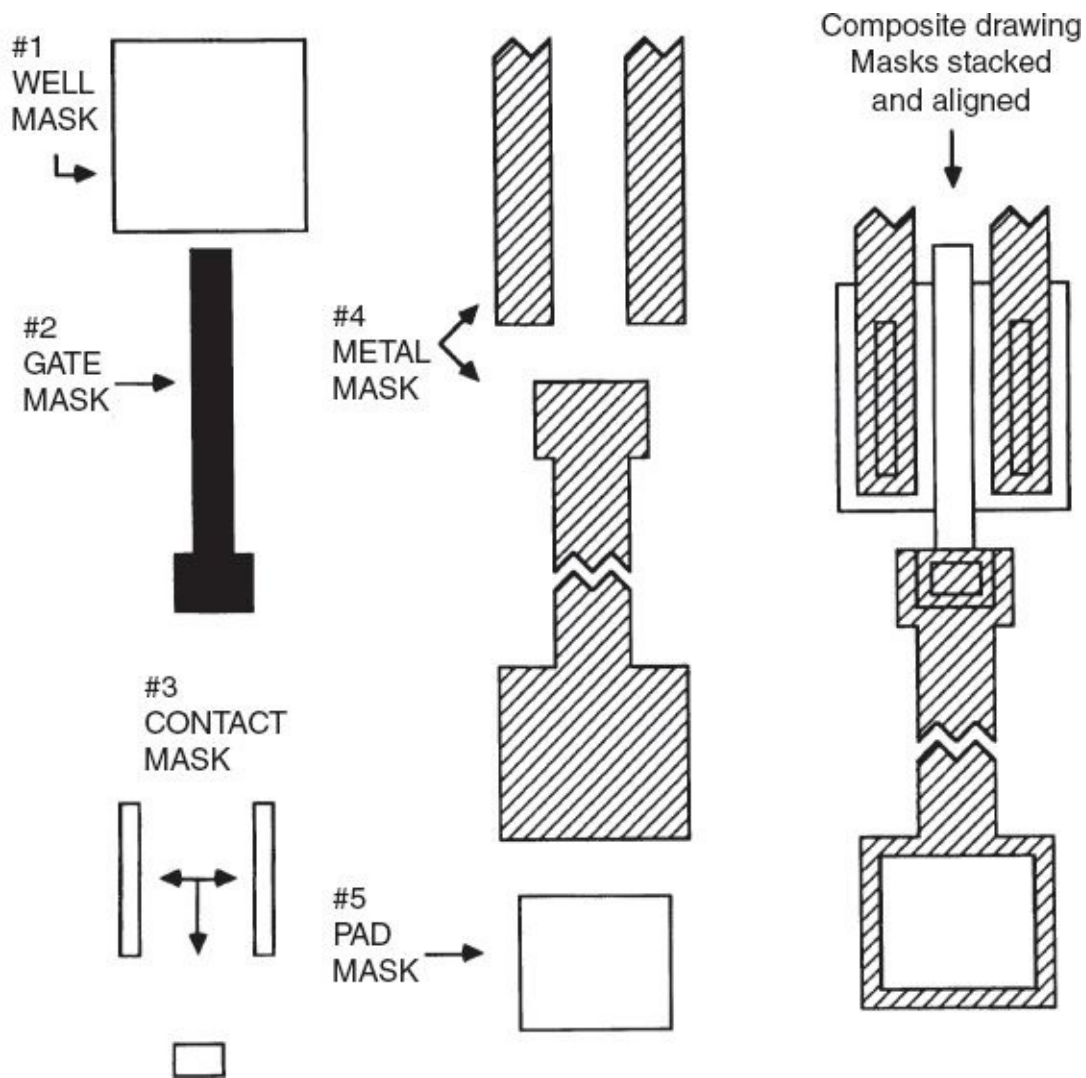


FIGURE 4.11 Composite and layer drawings for a five-mask silicon gate transistor.

Both buildings and semiconductor circuits are built one layer at a time. Therefore, it is necessary to separate the composite drawing into the layout for each individual layer in the circuit. [Figure 4.11](#) illustrates the composite and individual layer patterns for a simple silicon gate MOS transistor.

Each layer drawing is digitized (digitizing is the translation of the layer drawings to a digital database) and plotted on a computerized X-Y plotting table.

Reticle and Masks

The patterning process is used to create the required layer pattern and dimensions in and on the wafer surface. Getting the pattern from the digitized pattern to the wafer surface requires several steps. For the photo processes, there is an intermediate step called a reticle. A reticle is a “hard copy” of the individual drawing recreated in a thin layer of chrome deposited on a glass or quartz plate ([Fig. 4.12a](#)). It is created

using a photomasking process, the same as used on wafers using an electron beam (e-beam) generator for exposure of the photoresist. The reticle may be used directly in the patterning process or may be used to produce a photomask. A photomask is also a glass plate with a thin chrome layer on the surface. After production, it is covered with many copies of the circuit pattern ([Fig. 4.12b](#)). It is used to pattern a whole wafer surface in one pattern transfer. E-beam exposure systems skip the reticle or mask step and expose directly on the wafer surface. Reticle and mask-making processes are detailed in [Chap. 10](#).

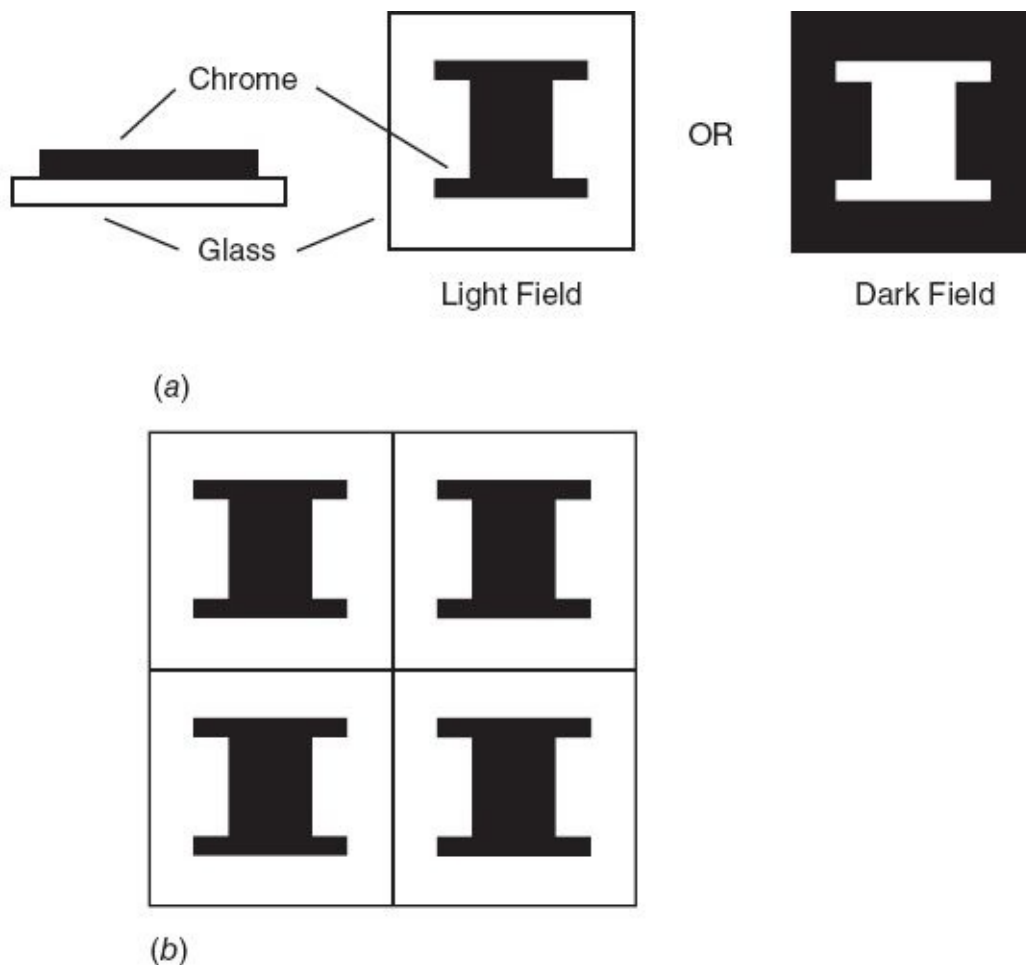


FIGURE 4.12 (a) Chrome on glass reticle and (b) photomask of the same pattern.

Reticles and masks are produced by an in-house department or purchased from outside vendors. Each circuit type has its own set of separate reticles or mask.

Basic Ten-Step Patterning Process

There are many individual patterning processes that address the proliferation of device structures, combinations of stacked layers in the devices, and the continued reduction in device dimensions. Yet there is a basic ten-step patterning process

illustrated in [Fig. 4.13](#). The details and variations are discussed in [Chaps. 8–10](#).

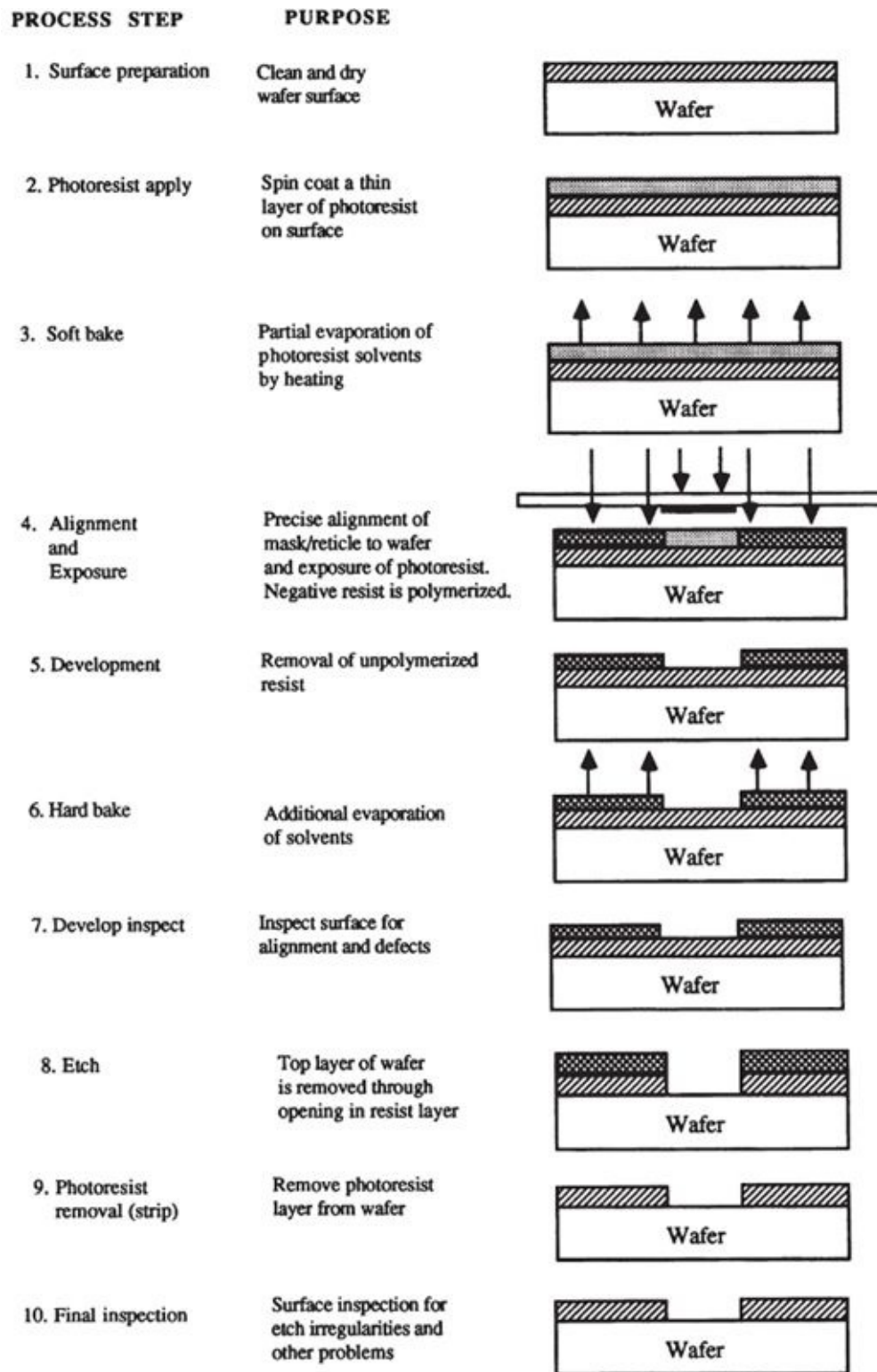


FIGURE 4.13 Ten-step photomasking process.

Doping

Doping is the process that puts specific amounts of electrically active dopants in the wafer surface through openings in the surface layers ([Fig. 4.14](#)). The two techniques are *ion implantation* and *thermal diffusion* which are detailed in [Chap. 11](#).

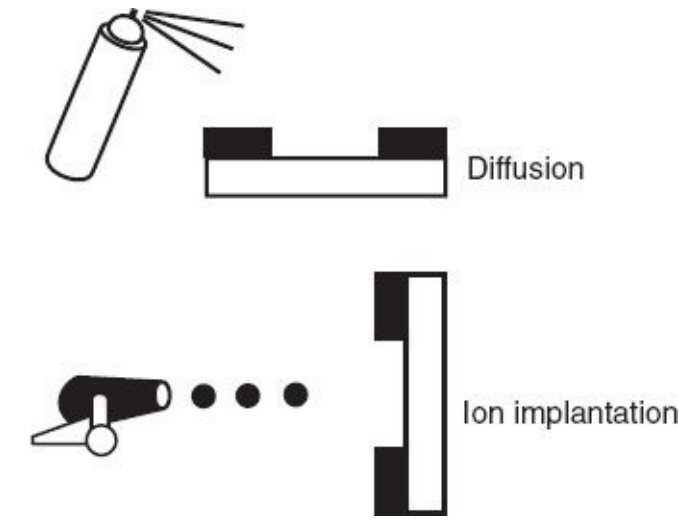


FIGURE 4.14 Doping techniques: diffusion and ion implantation.

Thermal diffusion is an older doping process using a chemical process that takes place when the wafer is heated to the vicinity of 1000°C and is exposed to vapors of the proper dopant. A common example of diffusion is the spreading of deodorant vapors into a room after being released from a pressurized can. Dopant atoms in the vapor state move into the exposed wafer surface through the chemical process of diffusion to form a thin layer in the wafer surface. In microchip applications, diffusion is also called solid-state diffusion, since the wafer material is a solid. Diffusion doping is a chemical process. Diffusion movement of dopants is governed by physical laws.

Ion-implantation doping is a physical process. Wafers are loaded in one end of an implanter and dopant sources (usually in gas form) in the other end. At the source end, the dopant atoms are ionized (given an electrical charge), accelerated to a high speed, and swept across the wafer surface. The momentum of the atoms carries them into the wafer surface, much like a ball shot from a cannon lodges in a wall.

The purpose of the doping operation is to create pockets in the wafer surface ([Fig. 4.15](#)) that are either rich in electrons (N-type) or rich in electrical holes (P-type). These pockets form the electrically active regions and N-P junctions required for operation of the transistors, diodes, capacitors, and resistors in the circuit.

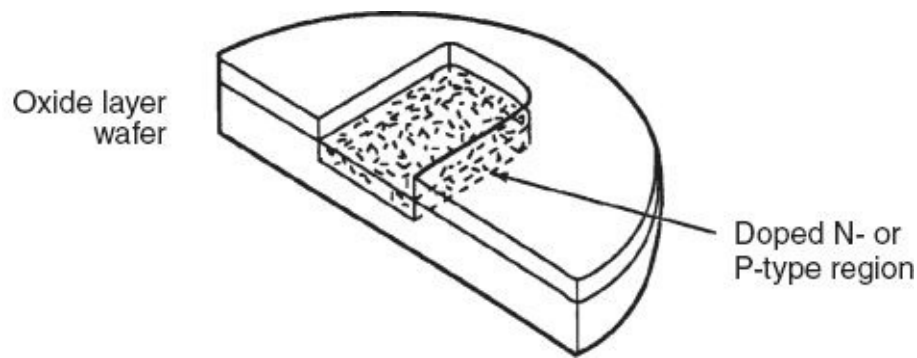


FIGURE 4.15 Formation of N- or P-type region in wafer surface.

Heat Treatments

Heat treatments are the operations in which the wafer is simply heated and cooled to achieve specific results ([Fig. 4.16](#)). In the heat-treatment operations ([Fig. 4.17](#)), no additional material is added or removed from the wafer. However, contaminants from the process may end up on or in the wafer.

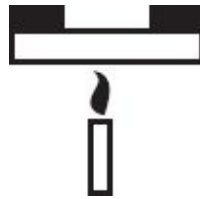


FIGURE 4.16 Heat treatments.

Operation	Heat Treatment
Patterning	Soft bake Hard bake Post-exposure bake (develop)
Doping	Post-ion implant anneal
Layering	Postmetal deposition and patterning anneal

FIGURE 4.17 Table of major heat treatments.

An important heat treatment takes place after ion implantation. The implantation of the dopant atoms causes a disruption of the wafer crystal structure that is repaired by a heat treatment, called *anneal*, at about 1000°C. Another heat treatment takes place after the conducting stripes of metal are formed on the wafer. These stripes carry the electrical current between the devices in the circuit. To ensure good electrical conduction, the metal is alloyed to the wafer surface by a heat treatment, which takes place at 450°C. A third important heat treatment is the heating of wafers

with photoresist layers to drive off solvents that interfere with accurate patterning.

Example Fabrication Process

The manufacture of a circuit starts with a polished wafer. The cross-section sequence in [Fig. 4.18](#) shows the basic operations required to form a simple MOS silicon-gate transistor structure. The following is an explanation of each operation in the fabrication process.

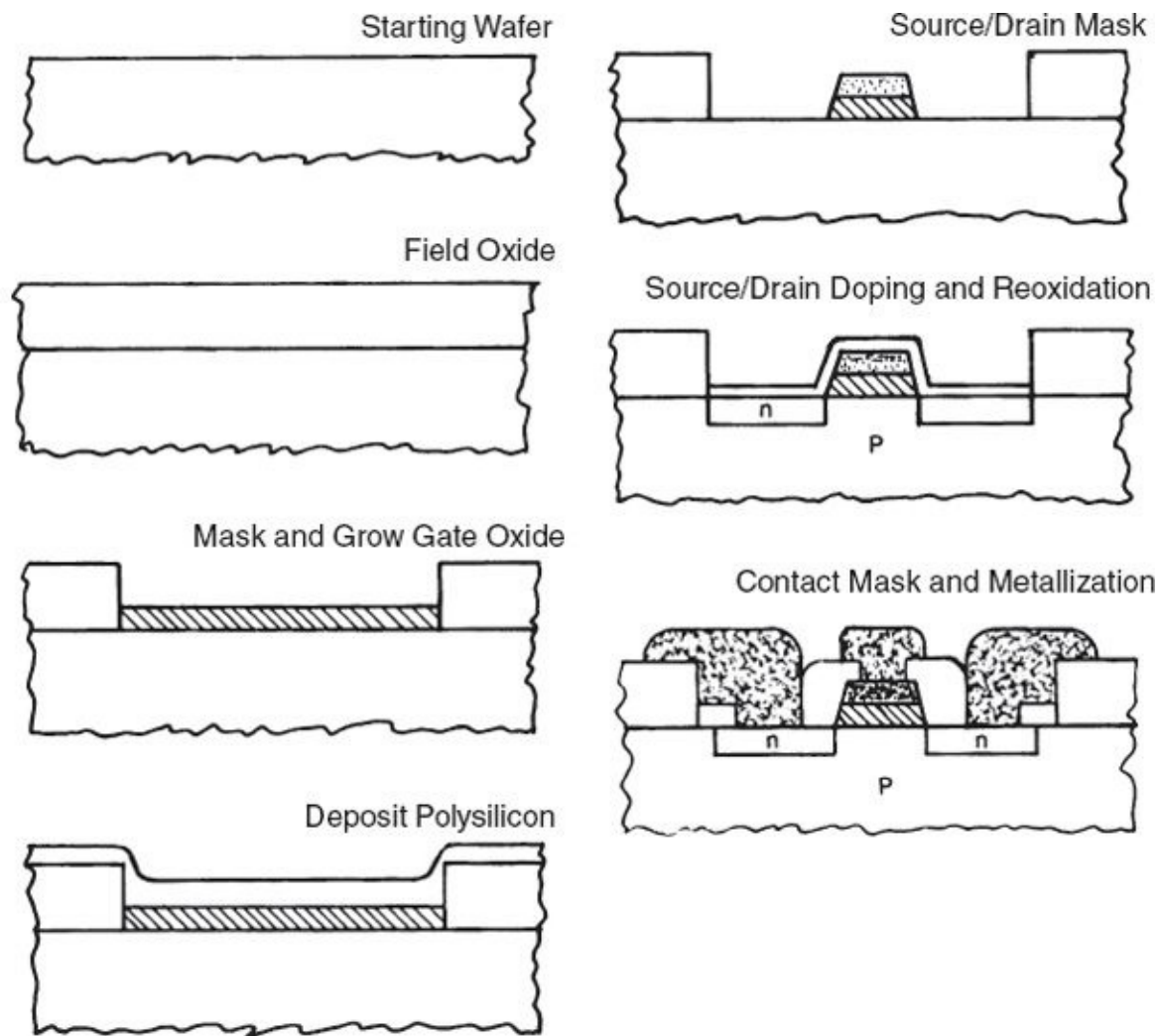


FIGURE 4.18 Silicon gate MOS process steps.

Step 1: Layering operation. The building starts with an oxidation of the wafer surface to form a thin protective layer and to serve as a doping barrier. This silicon dioxide layer is called the *field oxide*.

Step 2: Patterning operation. The patterning process leaves a hole in the field oxide that defines the location of the source, gate, and drain areas of the transistor.

Step 3: Layering operation. Next, the wafer goes to a silicon dioxide oxidation operation. A thin oxide is grown on the exposed silicon. It will serve as the gate oxide.

Step 4: Layering operation. In this step, another layering operation is used to deposit a layer of polycrystalline (poly) silicon. This layer will also become part of the gate structure.

Step 5: Patterning operation. Two openings are patterned in the oxide or polysilicon layer to define the source and drain areas of the transistor.

Step 6: Doping operation. A doping operation is used to create an N-type pocket in the source and drain areas.

Step 7: Layering operation. Another oxidation or layering process is used to grow a layer of silicon dioxide over the source or drain areas.

Step 8: Patterning operation. Holes, called contact holes, are patterned in the source, gate, and drain areas.

Step 9: Layering operation. A thin layer of conducting metal, usually an aluminum alloy, is deposited over the entire wafer.

Step 10: Patterning operation. After deposition, the wafer goes back to the patterning area where portions of the metallization layer are removed from the chip area and the scribe lines. The remaining portions connect all the parts of the surface components to each other in the exact pattern required by the circuit design.

Step 11: Heat treatment operation. Following the metal patterning step, the wafer goes through a heating process in a nitrogen-gas atmosphere. The purpose of the step is to “alloy” the metal to the exposed source and drain regions and the gate region to ensure good electrical contact.

Step 12: Layering operation. The final layer of this device is a protective layer known variously as a *scratch* or *passivation layer*. Its purpose is to protect the components on the chip surface during the testing and packaging processes, and during use.

Step 13: Patterning operation. The last step in the sequence is a patterning process that removes portions of the scratch protection layer over the metallization terminal pads on the periphery of the chip. This step is known as the *pad mask*.

The 13-step process illustrates how the four basic fabrication operations are used to build a particular transistor structure. The other components (diodes, resistors, and capacitors) required for the circuit are formed in other areas of the circuit as the transistors are being formed. For example, in this sequence, resistor patterns are put on the wafer at the same time as the source or drain pattern for the transistor. The subsequent doping operation creates the source or drain *and* the resistors. Other transistor types, such as the bipolar and silicon gate MOS, are formed by the same basic four operations, but using different materials and in different sequences.

In general, the circuit components are formed in the first series of fabrication

operations and referred to as the *front end of the line* (FEOL). In the later series of operations, the various metallization layers that connect the circuit components are added to the wafer surface. These operations are called the *back end of the line* (BEOL).

Modern chip structures are many times more complicated than the simple process just described. They have multiple layers and pockets of dopants, numbers of layers added to the surface, including multiple layers of conductors interspersed with dielectrics ([Fig. 4.19](#)).

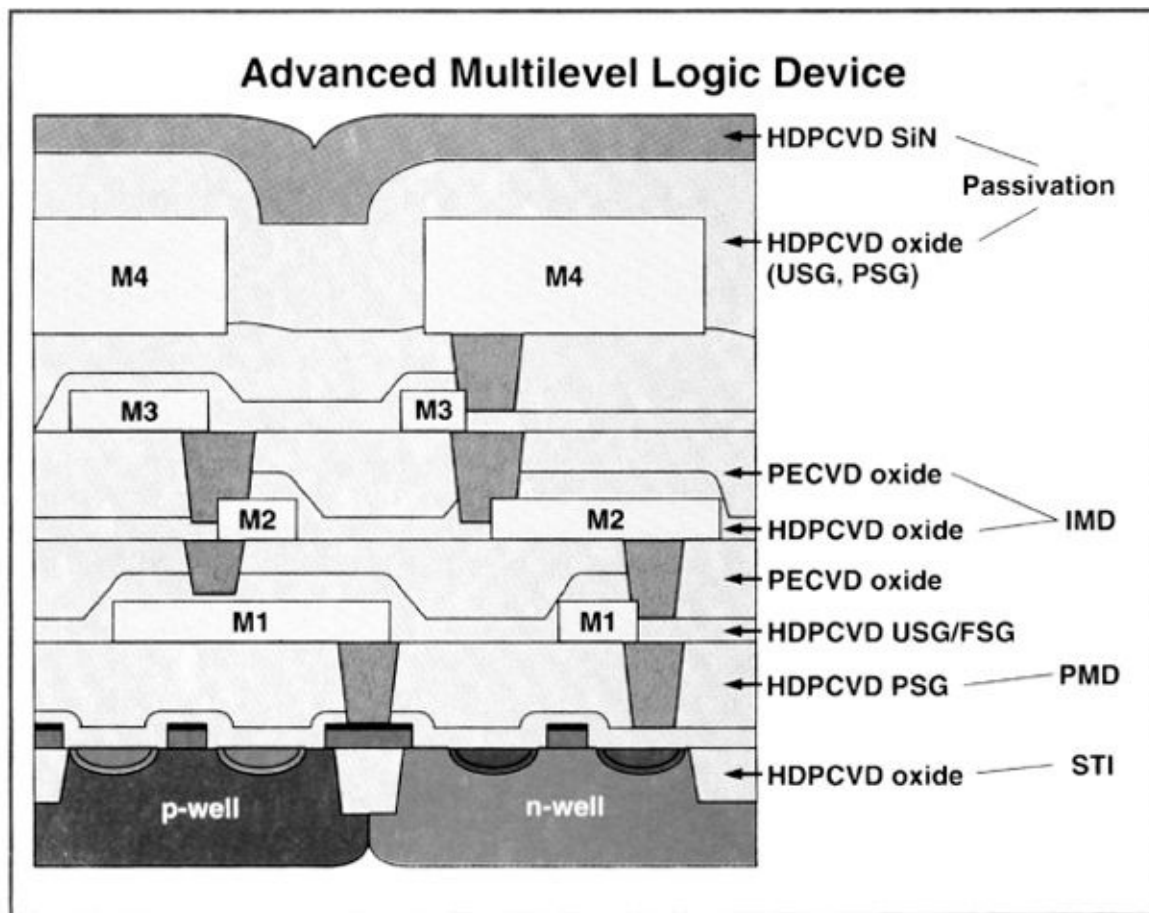


FIGURE 4.19 Modern chip structure.

Achieving these complicated structures requires many processes. And each process, in turn, requires a number of steps and substeps. A speculative process for a 64-Gb CMOS device might require 180 major steps, 52 clean and/or strips, and up to 28 masks.⁴ Yet all of the major steps are one of the four basic operations.

By the time the industry reaches circuits with gate widths of several atoms and stacks of metal on top of the circuit, the number of process steps will be 500 or more.

Wafer Sort

Following the wafer-fabrication process comes a very important testing step: wafer sort. This test is the report card on the fabrication process. During the test, each chip is electrically tested for electrical performance and circuit functioning. *Wafer sort* is also known as *die sort* or *electrical sort*.

For the test, the wafer is mounted on a vacuum chuck and aligned to thin electrical probes that contact each of the bonding pads on the chip (Fig. 4.20). The probes are connected to power supplies that test the circuit and record the results. The number, sequence, and type of tests are directed by a computer program. Wafer probers are automated so that, after the probes are aligned to the first chip (manually or with an automatic vision system), subsequent testing proceeds without operator assistance.

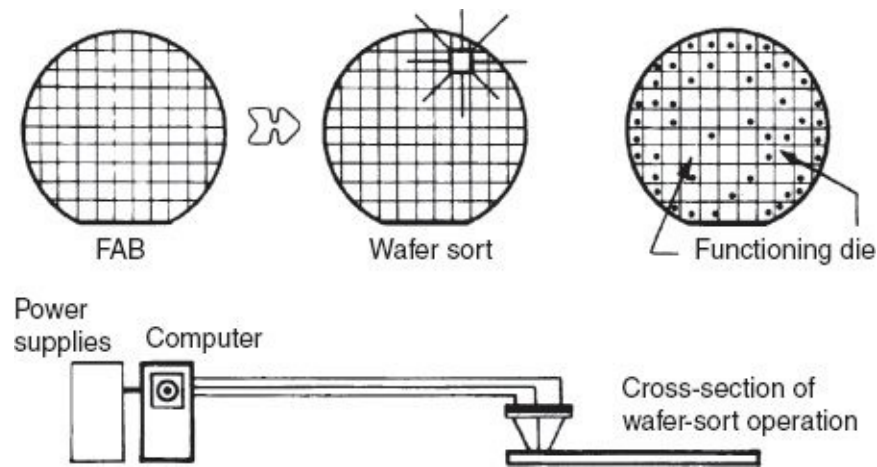


FIGURE 4.20 Wafer sort.

The goal of the test is threefold. First is the identification of working chips before they go into the packaging operation. Second is characterization of the device or circuit electrical parameters. Engineers need to track the parameter distributions to maintain process quality levels. The third goal is an accounting of the working and nonworking chips to give fab personnel feedback on overall performance (yield). The location of the working and nonworking chips is logged into a wafer map in the computer. Older technologies deposit a drop of ink on the *nonworking* chips.

Wafer sort is one of the principal yield calculations in the chip-production process. It also gets more expensive as the chips get larger and denser.² These chips require more time to probe, power supplies, wafer handling mechanics, and computer systems have to be more sophisticated to perform the tests and track results. The vision systems must also evolve in sophistication (and cost) with expanding die size. Cutting the chip test time is also a challenge. Chip designers are being asked to include test modes for memory arrays. Test designers are exploring ways to streamline test sequences using stripped-down tests once the chip is fully

characterized, performing scan tests of the circuit, and parallel testing different circuit parts. The details of wafer parameter yields are addressed in [Chap. 6](#).

Packaging

Most wafers continue on to stage four, called *packaging* ([Fig. 4.21](#)). The wafers are transferred to a packaging area on the same site or to a remote location. Many semiconductor producers package their chips in offshore facilities. ([Chap. 18](#) details the packaging process.) In this stage, the wafers are separated into the individual chips, and the working chips are incorporated into a protective package. Some chips are directly incorporated into electrical systems without a package.

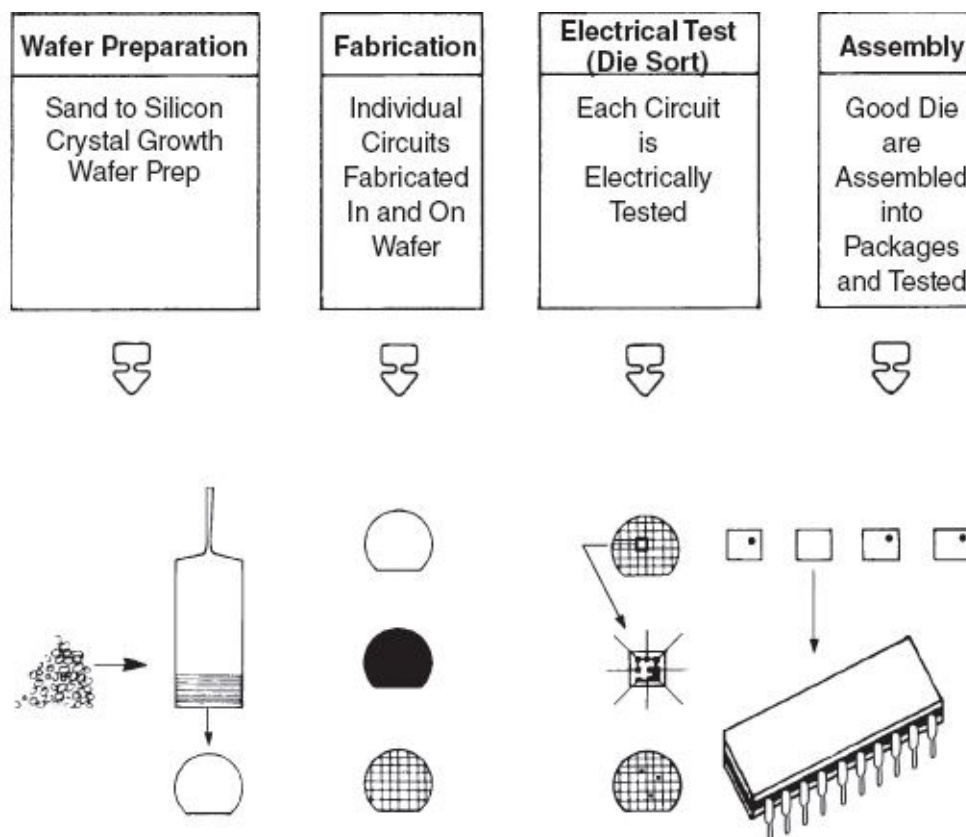


FIGURE 4.21 Integrated circuit manufacturing sequence.

Summary

The semiconductor-microchip-fabrication process is long, complicated, and includes many variations, depending on the type of product, level of integration, feature size, and other factors. Understanding this particular process is easier by separating it into the four stages. Wafer fabrication is further understood by identifying the four basic operations performed on the wafer. Several simple

processes have been used to illustrate the basic wafer-fabrication operations. Actual processes used are addressed in the process chapters and in [Chaps. 16](#) and [17](#). Industry drivers and manufacturing trends are explained in [Chap. 15](#).

Review Topics

Upon completion of this chapter, you should be able to: 1. Identify and explain the four basic wafer operations.

2. Identify the parts of a wafer.
3. Draw a flow diagram of the circuit-design process.
4. Explain the definition and use of a composite drawing and mask set.
5. Draw cross-sections showing the doping sequence of basic operations.
6. Draw cross-sections showing the metallization sequence of basic operations.
7. Draw cross-sections showing the passivation sequence of basic operations.
8. Identify the parts of an integrated circuit chip.

References

1. Kopp, R., *Kopp Semiconductor Engineering*, 1996.
2. Iscoff, R., "VLSI Testing: The Stakes Get Higher," *Semiconductor International*, Sep. 1993:58.

CHAPTER 5

Contamination Control

Introduction

The effects of contamination on device processing, device performance, and device reliability are explained and identified along with the types and sources of contamination found in a fabrication area. Cleanroom layouts, major contamination-control procedures, and wafer-surface-cleaning techniques are explained.

One of the first problems to plague the infant microchip fabrication efforts was contamination. The industry started with a cleanroom technology developed by the space industry. However, these techniques proved inadequate for large-scale manufacturing of chips. Cleanroom technology has had to keep pace with chip design and density evolution. The ability of the industry to grow has been dependent on solving contamination problems presented by each generation of chips. Yesterday's minor problems become the killer defects of today's chips.

The Problem

Semiconductor devices are very vulnerable to many types of contaminants. They fall into five major classes:

1. Particles
2. Metallic ions
3. Chemicals
4. Bacteria
5. Airborne molecular contaminants (AMCs)

Particles

Semiconductor devices, especially dense integrated circuits, are vulnerable to all kinds of contamination. The sensitivity is due to the small feature sizes and the thinness of deposited layers on the wafer surface. These dimensions are down to the submicron range. A micron or micrometer (μm) is very small. Current device dimensions (feature size) are now in the nanometer territory (nm). A way to envision a micron is that a human hair is about $100\ \mu\text{m}$ in diameter ([Fig. 5.1](#)). The small physical dimensions of the devices make them very vulnerable to particulate contamination in the air coming from workers, generated by the equipment, and present in processing chemicals ([Fig. 5.2](#)). As the feature size and films become smaller ([Fig. 5.3](#)), the allowable particle size must be controlled to smaller dimensions.

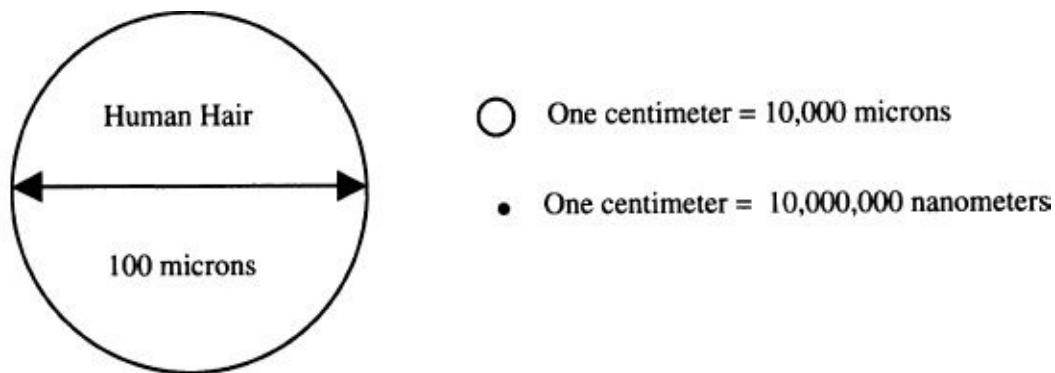


FIGURE 5.1 Relative sizes.

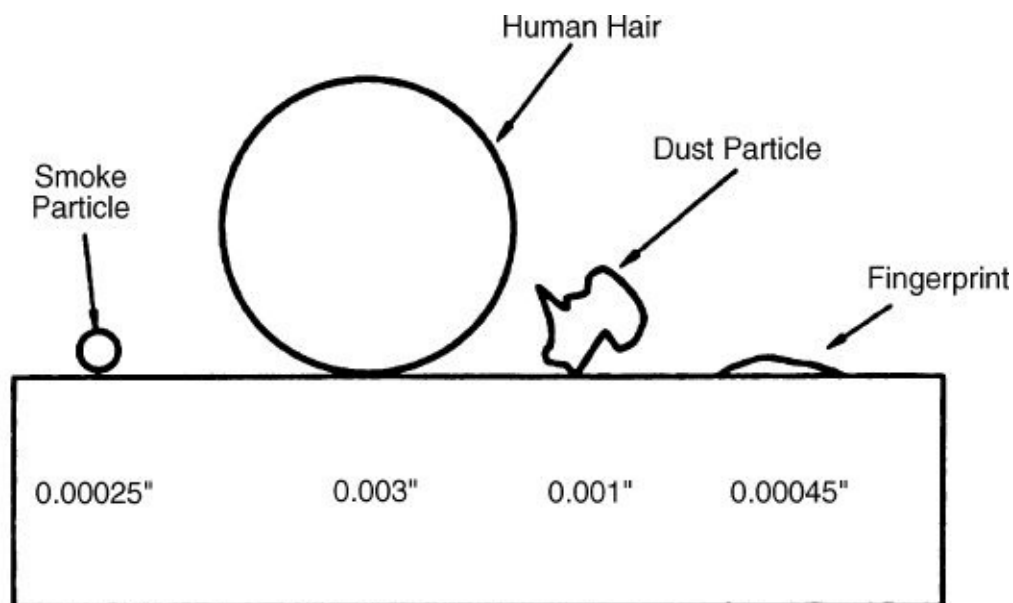


FIGURE 5.2 Relative size of contamination. (Source: Hybrid Microcircuit Technology Handbook, 1988:281.)

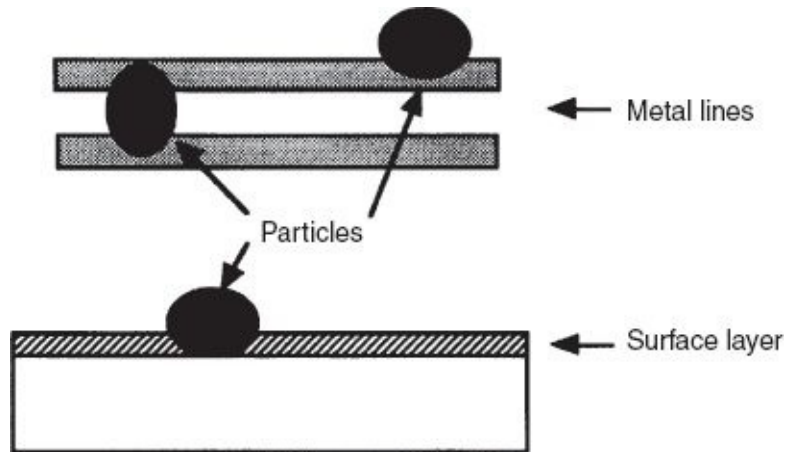


FIGURE 5.3 Relative size of airborne particles and wafer dimensions.

A rule of thumb is that the particle size must be $\frac{1}{2}$ the first metal layer *half pitch*.¹ A half pitch is $\frac{1}{2}$ the distance from an adjacent metal strip and space. Particles that locate in a critical part of the device and destroy its functioning are called *killer defects*. Killer defects also include crystal defects and other process-induced problems. On any wafer, there are a number of particles. Some are the killer variety, but others locate in less sensitive areas and do no harm. The Yield Enhancement chapter of the 2011 Edition of the *International Technology Roadmap for Semiconductor* (ITRS) identifies defect correlation with yield and development of more sensitive defect and contamination detection tools as essential to yield enhancement for future technology nodes.

Metallic Ions

In [Chapter 2](#), it was established that semiconductor devices require controlled resistivity in the wafers and in the doped N- and P-regions, and precise N-P junctions. These three properties are achieved by the purposeful introduction of specific dopants into the crystal and into the wafer. These desired effects are achieved with very small amounts of the dopants. Unfortunately, it takes only a small amount of certain electrically active contaminants in the wafer to alter device electrical characteristics, changing performance and reliability parameters.

The contaminants causing these types of problems are known as *mobile ionic contaminants* (MICs). They are atoms of metals that exist in the material in an ionic form (as they carry an electric charge). Furthermore, these metallic ions are highly mobile in semiconductor materials. This mobility means that the metallic ions can move inside the device, even after passing electrical testing and shipping, causing the device to fail. Unfortunately, the metals ([Fig. 5.4](#)) that cause these problems in silicon devices are present in most chemicals. On a wafer, MIC contamination has to be in the 10^{10} atoms/cm² range or less.²

Chemical grade	Customer application width in microns	Max. metallic impurities specification per element
SS-ULSI	<0.13	0.1 ppb / 100 ppt
S-ULSI	0.13 - 0.8	1 ppb
ULSI	>0.8	10 ppb
VLSI	>1.0	100 ppb
MOS	>1.0	1000 ppb

FIGURE 5.4 Metallic impurities specifications.

Sodium is the most prevalent mobile ionic contaminant in most untreated chemicals and is the most mobile in silicon. Consequently, control of sodium is a prime goal in silicon processing. The MIC problem is most serious in metal oxide silicon (MOS) devices, which has led specialty chemical suppliers to develop MOS or low-sodium-grade chemicals. And ultrapure water manufacturing also requires reduction of MICs.

Chemicals

The third major contaminant in semiconductor process areas is unwanted chemicals. Process chemicals and process water can be contaminated with trace chemicals that interfere with the wafer processing. They may result in unwanted etching of the surface, creation of compounds that cannot be removed from the device, or causing nonuniform processes. Chlorine is such a contaminant and is rigorously controlled in process chemicals.

Bacteria

Bacteria is the fourth major contaminant class. Bacteria are organisms that grow in water systems and on surfaces that are not cleaned regularly. Bacteria, once on the device, act as particulate contamination or may contribute unwanted metallic ions to the device surface.

Airborne Molecular Contaminants

Airborne molecular contaminants (AMCs) are fugitive molecules that escape from the process tools or the chemical delivery systems, or are carried into the fabrication area on materials or personnel. Transfer of wafers from one process tool to another can carry hitchhiking molecules to the next tool. AMCs include all of the gases, dopants, and process chemicals used in the fabrication area. These can be oxygen, moisture, organics, acids, bases, and others.³

Their harm is highest in processes that involve delicate chemical reactions, such as the exposure of photoresist in the patterning operations. Other problems include the shifting of etch rates and unwanted dopants that shift device electrical parameters and change the wetting characteristics of etchants leading to incomplete etching.⁴ Detecting and controlling AMCs is a must as automation introduces more equipment and environments into the fabrication process. One source of concern identified in the 2011 *International Technology Roadmap for Semiconductor* (ITRS), Yield Enhancement chapter is front opening universal pods (FOUPs). The plastic material in these wafer containers are a source of AMC outgassing. Also the wafers can outgas materials and by-products of previous process steps.

Contamination-Caused Problems

The five types of contaminants affect the processing and devices in three specific performance areas include:

1. Device processing yield
2. Device performance
3. Device reliability

Device Processing Yield

Device processing in a contaminated environment can cause a multitude of problems. Contamination may change the dimensions of device parts, change the cleanliness of the surfaces, and/or cause pitted layers. Within the fabrication process, there are many quality checks and inspections specifically designed to detect contaminated wafers. Contamination-caused defects contribute to rejection of wafers in the wafer-fabrication process, reducing the overall yield (see [Chap. 6](#)).

Device Performance

A more serious problem is related to small pieces of contamination that may escape the in-process quality checks. Also unwanted chemicals and AMCs in the process steps may alter device dimensions or material quality. High levels of mobile ionic contaminants in the wafer can change the electrical performance of the devices. These problems usually show up at an electrical test (called *wafer* or *die sort*) that checks each chip after the wafer-fabrication process (see [Chap. 6](#)).

Device Reliability

Loss of device reliability is the most insidious of the contamination failures. Small amounts of metallic contaminants can get into the wafer during processing and not be detected during normal device testing. However, in the field, these contaminants can travel inside the device and end up in electrically sensitive areas, causing failure. This failure mode is a primary concern of the space and defense industries.

In the rest of this chapter, the sources, nature, and control of the types of contamination that affect semiconductor devices are identified. With the advent of large scale integration (LSI) level circuits in the 1970s, the control of contamination became essential to the industry. Since that time, a great deal of knowledge about and control of contamination has been learned. Contamination control is now a discipline of its own and is one of the critical technologies that have to be mastered to profitably produce solid-state devices.

Contamination control issues addressed in this chapter apply to wafer-fabrication areas, mask-making areas, some chip packaging areas, and areas in which semiconductor equipment and materials are manufactured.

Contamination Sources

General Sources

Contamination in a cleanroom is defined as anything that interferes with the production of the product and/or its performance. The stringent requirements of solid-state devices define levels of cleanliness that far exceed those of almost any other industry. Literally everything that comes in contact with the product during manufacture is a potential source of contamination. The major contamination sources are:

1. Air
2. The production facility
3. Cleanroom personnel
4. Process water
5. Process chemicals
6. Process gases
7. Static charge
8. Process equipment

Each source produces specific types and levels of contamination and requires special controls to render it acceptable in the cleanroom.

Air

Normal air is so laden with contaminants that it must be treated before entering a cleanroom. A major problem is airborne particles, referred to as *particulates* or *aerosols*. Normal air contains copious amounts of small dust and particles, as illustrated in [Fig. 5.5](#). A major problem with small particles (called *aerosols*) is that they “float” and remain in the air for long periods of time. Air cleanliness levels in cleanrooms are identified by the particulate diameters and their density in the air.

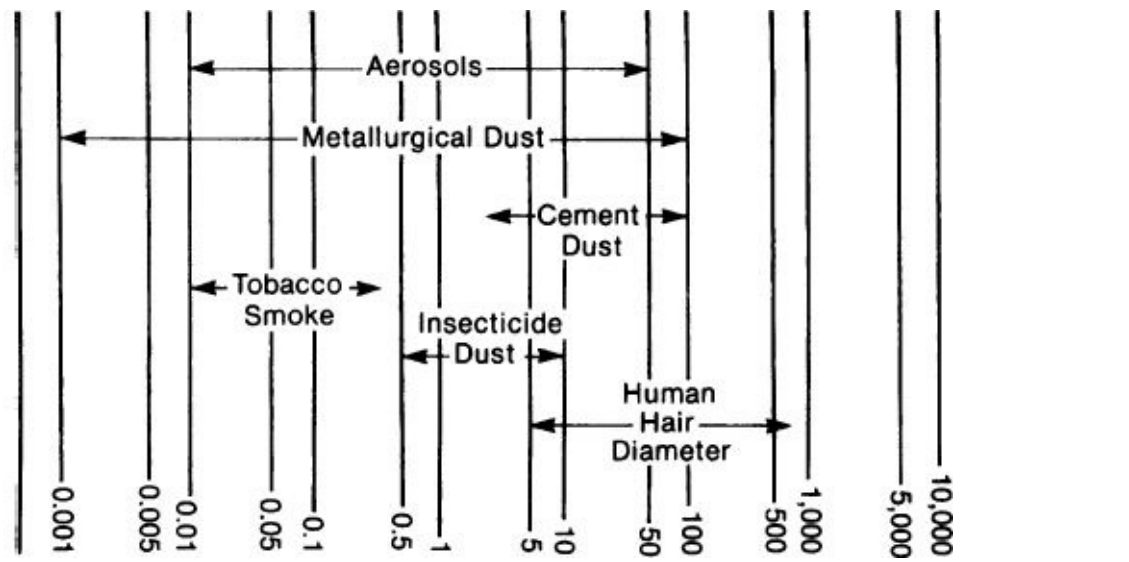
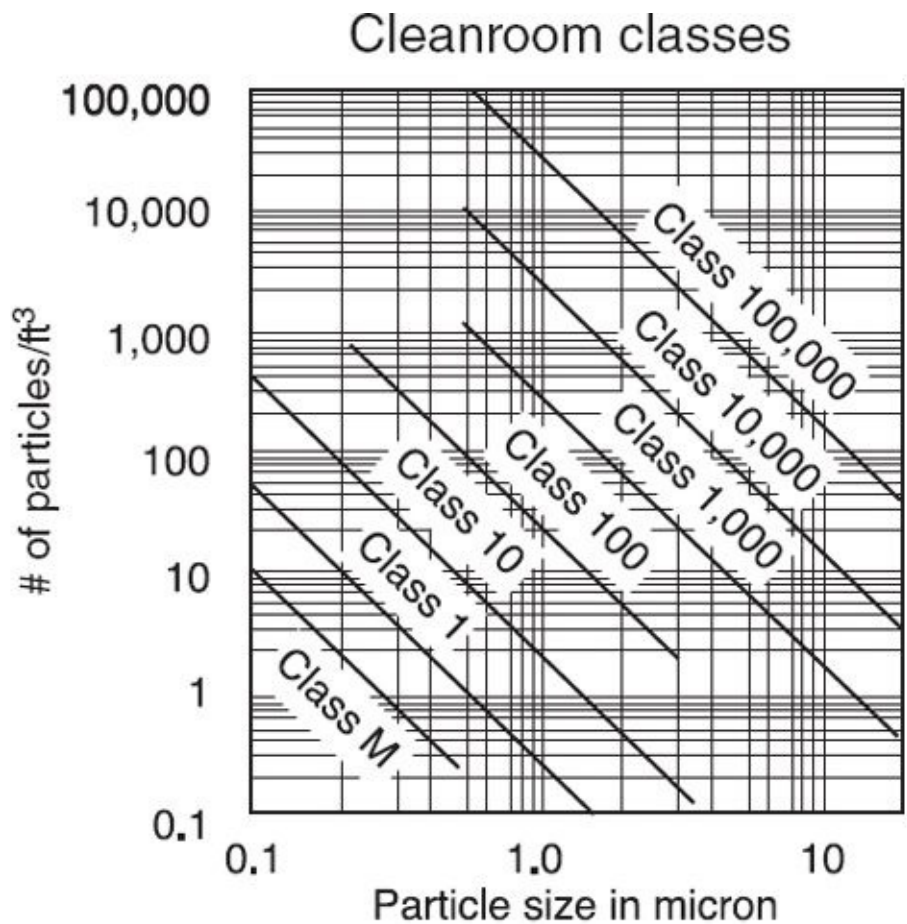


FIGURE 5.5 Relative size of airborne particulates (in microns).

Air quality is designated by the *class number* of the air in the area as originally defined in Federal Standard 209E.⁵ An International Standard (ISO) has replaced 209E and added additional classes. This standard designates air quality in the two categories of particle size and density. The class number of an area is defined as the number of particles of different sizes in a cubic foot of air. The air in a typical city, filled with smoke, smog, and fumes, can contain up to 5 million particles per cubic foot, which is a class number of 5 million. Advancing chip sensitivity has identified smaller and smaller tolerable particle sizes for each generation of chip feature size.⁶

[Figure 5.6](#) shows the relationship between particulate diameter and density as defined by the ISO/Federal Standard 209E. Class M-1 reflects the higher cleanliness standards required for 300-or 450-mm fab production areas. [Figure 5.7](#) lists the class numbers and associated particle size for various environments. Federal Standard 209E specifies cleanliness levels down to class 1 levels. While 209E defines class numbers at 0.5- μm particle size, successful wafer-fabrication processing requires tighter controls. Engineers strive to achieve reduction of 0.1-

μm particles in class 10 and class 1 environments. Specifications proposed by Semetech call for process areas at class 0.1 and air levels of class 0.01.²



Definition of Airborne Particulate Cleanliness
Class per Fed. Std. 209E

Class	Particles/ft ³				
	0.1 um	0.2 um	0.3 um	0.5 um	5 um
M-1	9.8	2.12	0.865	0.28	
1	35	7.5	3	1	
10	350	75	30	10	
100		750	300	100	
1000				1000	7
10,000				1000	70

FIGURE 5.6 Air cleanliness class standards—209E/ISO.

Environment	Class size	Maximum particle size (micron)
Projected-256 Mbit	0.01	<<0.1
Mini-environment	0.1	<0.1
ULSI Fab	1	0.1
VLSI Fab	10	0.3
VLF Station	100	0.5
Assembly Area	1000–10,000	0.5
House room	1000,000	
Outdoors	>500,000	

FIGURE 5.7 Typical class numbers for various environments.

Clean Air Strategies

The design of a cleanroom is integral to its ability to produce contamination-free wafers. A major consideration in the design is the maintenance of clean air in the process areas. Automation is also an important cleanroom strategy for the reduction of contamination. This issue is explored in the equipment section and in [Chap. 15](#). Four distinct room design strategies are used:

1. Ballroom
2. Tunnel design
3. Minienvironments
4. Wafer isolation technology (WIT)

Cleanroom Workstation Strategy

The semiconductor industry adopted cleanroom techniques first developed by NASA to assemble space vehicles and satellites. However, the small rooms adequate for satellite assembly could not be maintained when expanded to larger fabrication areas with more production workers. This problem was addressed with a cleanroom workstation strategy. Outside the workstations, the wafers were stored and moved in covered boxes.

The initial fabrication area consisted of a large room referred to as a *ballroom* design. Workstations (called *hoods*) were arranged in rows so that the wafers could move sequentially through the process, never being exposed to dirty air. The filters in these clean hoods are either high-efficiency particulate attenuation (HEPA) filters or, for more demanding applications, ultra-low-particle (ULPA) filters. These filters are constructed of fragile fibers with many small holes and are folded into the filter holder in an accordion design ([Fig. 5.8](#)). The high density of small holes and large area of the filter medium allow the passage of large volumes of air at low velocity. The low velocity contributes to the cleanliness of the hood by not causing air currents. The low velocity is also necessary for operator comfort. A typical airflow is 90 to 100 ft/min. HEPA and ULPA filters achieve filtering efficiencies as high as 99.9999+ percent at 0.12 μm or lower particle size.⁸

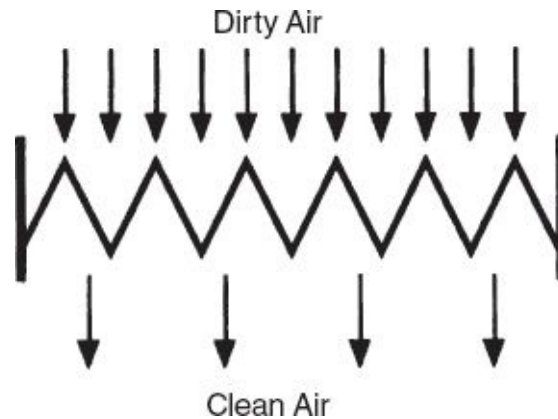


FIGURE 5.8 HEPA filter.

Workstations come in two varieties. A clean hood ([Fig. 5.9](#)) has a HEPA/ULPA filter mounted on the top. Air is drawn from the room through a prefilter by a fan and forced through the HEPA filter. The air leaves the filter in a laminar pattern and, at the work surface, it turns and exits the hood. A shield directs the exiting air over the wafers in the hood. The formal name for the workstation is a *vertical laminar flow (VLF) station*. The term VLF is derived from the laminar nature of the airflow.

For chemical processes, the clean hoods are connected to the factory exhaust system contain the vapors and to a drain system to remove spent liquid chemicals ([Fig. 5.10](#)).

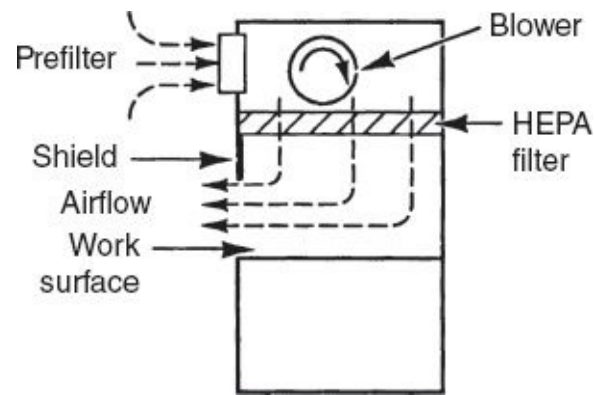


FIGURE 5.9 Cross-section of VLF process section.

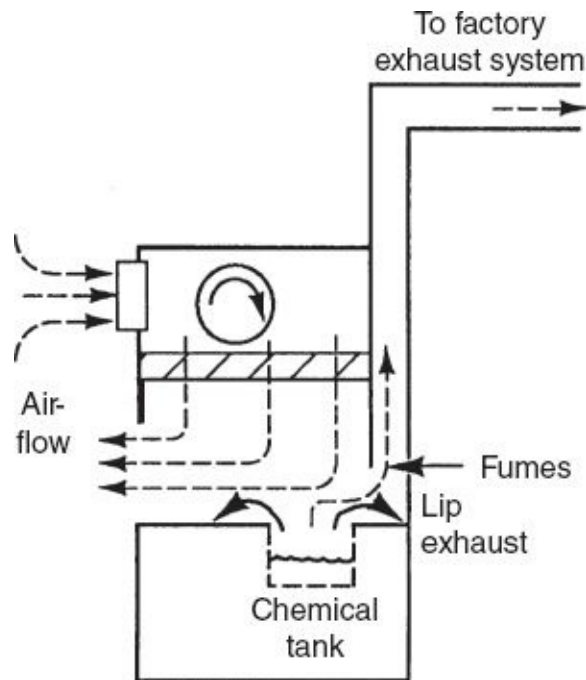


FIGURE 5.10 Cross-section of a VLF exhaust process section.

Both types of stations keep the wafers clean in two ways. First is the filtered air inside the hood. The second cleaning action is the slight positive pressure built up in the station. This pressure prevents airborne dirt from operators and from the aisle areas from entering the hood.

The basic clean-work station design also applies to equipment in modern fabs. Individual tools must be fitted with a VLF-or HLF-designed load section to keep the wafers clean during loading and unloading steps ([Chap. 15](#)).

Tunnel or Bay Concept

As more critical particulate control became necessary, it was noted that the VLF hood approach had several drawbacks. Chief among them was the vulnerability to contamination from the many personnel moving about in the room. People entering and exiting the fabrication area had the potential of contaminating all the process stations in the area.

This particular problem is solved by dividing the fabrication area into separate tunnels or bays (Fig. 5.11). Instead of the individual VLF hoods, filters are built into the ceilings and serve the same purpose. The wafers are kept clean by the filtered air from the ceiling filters and are less vulnerable to personnel-generated contamination because of fewer workers in the immediate vicinity. On the downside, tunnel arrangements cost more to construct and are less versatile when the process changes.⁹

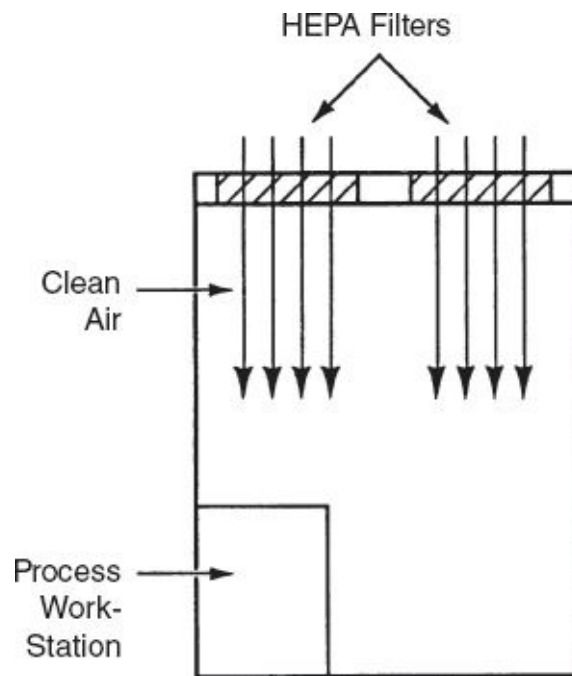
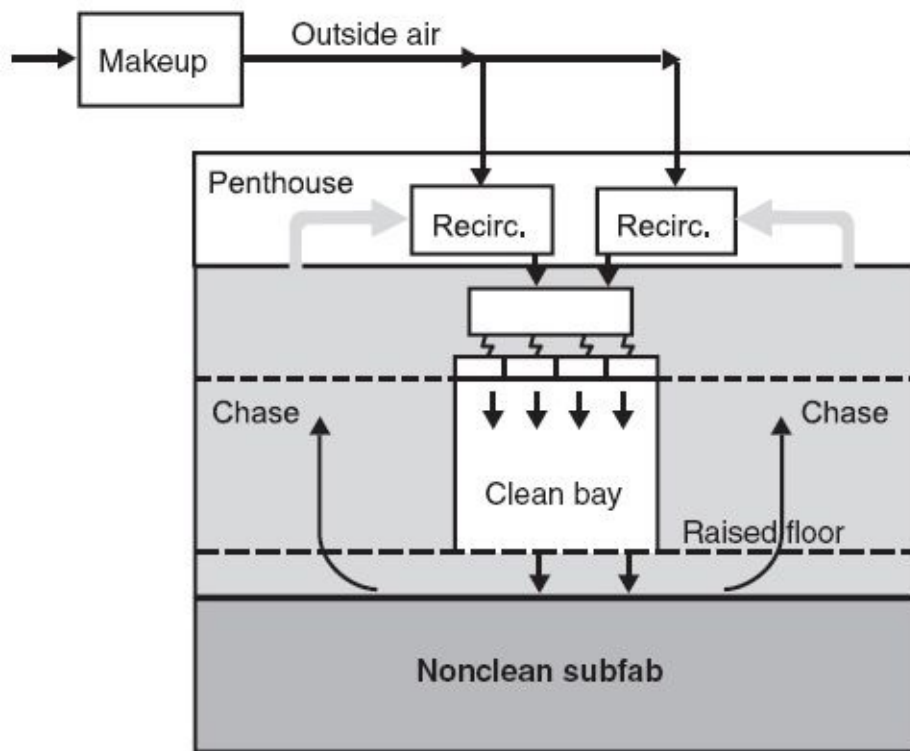


FIGURE 5.11 Cross-section of a traditional laminar flow cleanroom. (*Jamison Traditional Cleanroom.*)

The trend of equipment and facility design has been to isolate the wafer from contamination sources (Fig. 5.12). VLF hoods isolated the wafer from the room air, and tunnels isolate the wafer from excessive personnel exposure. The advent of CMOS ICs increased the number of process steps and the need to include more process stations in the cleanroom.

Traditional cleanroom



Bay and chase with nonclean subfab

FIGURE 5.12 Modern fab cross-section.

These larger rooms (and tunnels) bring with them the potential of contamination from the sheer volume of the air and the increased number of operators.

Micro-and Mini-Environments

Projections in the mid-1980s showed increasing cleanroom costs with diminishing returns on effectiveness. The concept of isolating the wafer in as small an environment as possible became the new direction. This concept was already in place with steppers and other process tools that had built in clean *microenvironments* for wafer loading and unloading ([Fig. 5.13](#)).

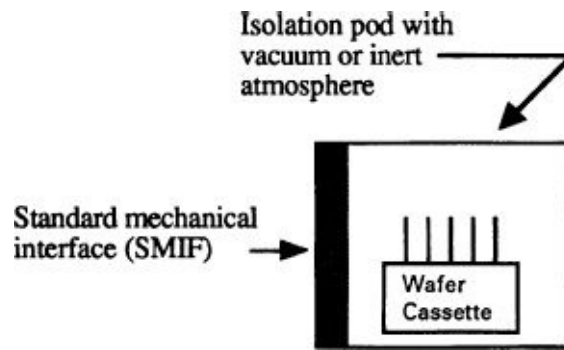


FIGURE 5.13 Wafer transfer SMIF box or front-loading universal pod.

The challenge was to string together a series of mini-environments such that the wafer was never exposed to the room air. Hewlett Packard developed a critical link with the invention, in the mid 1980s, of the *standard mechanical interface* (SMIF).¹⁰ With SMIF, the traditional wafer box is replaced with a wafer enclosure (*mini-environment*) that can be pressurized with air or nitrogen to keep out room air. This approach took on the general name of *wafer isolation technology* (WIT) or a mini-environment system. There are three parts to the system: the wafer box or pod for transporting the wafers, the isolated microenvironment at the tool, and a mechanism for extracting and loading the wafers. The wafer pods, called *SMIF* boxes evolved into FOUPs (front opening universal pod). These pods hold a batch of wafers keeping them isolated from the fab environment. However, contamination can come from wafer outgassing of chemicals picked up in previous process operations. FOUPs feature a mechanical interface that allows direct connection to the microenvironment of the process tool (Fig. 5.13). Wafers may be loaded directly from and to the pod onto the tool wafer system by dedicated handlers. Another approach is to move the cassette from the pod to the tool-wafer-handling system with a robot. Minienvironments offer the advantage of greater temperature and humidity control.

The WIT or mini-environment strategy includes the benefit of upgrading existing fabs along with other benefits (Fig. 5.14). Yield losses from contamination are lowered. This critical benefit can be delivered at lower facility construction and operating costs. WIT allows keeping the aisle air cleanliness at a lower cleanliness level, which reduces construction and operating costs. With the wafer isolated, there is less pressures on operator clothing, procedures, and constraints. However, the advent of larger-diameter wafers has driven up the weight of a pod of wafers to the point that they are too heavy for operators and too expensive to risk dropping. This situation requires robot handling, which increases cost and complexity. Minienvironment layouts must include storage for wafers (in pods) waiting for process tool availability. Current technology calls for storage units, called *stockers*, which hold the waiting pods and wafers. Layouts may include a central storage system with or without buffer storage at each tool. Isolation systems also zone off

process areas (photolithography, CVD & Doping, CMP & Wet Etch, specialty metals) from each other on the floor and in the exhaust system to prevent chemical cross-contamination.

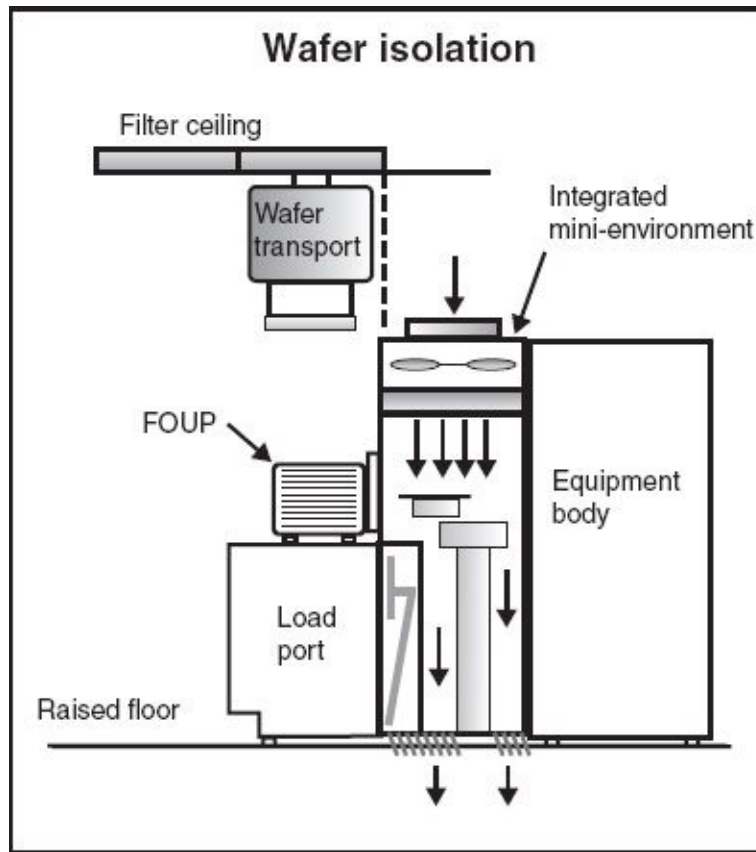


FIGURE 5.14 Wafer isolation technology.

Temperature, Humidity, and Smog

In addition to control of particulates, the air parameters of temperature, humidity, and smog must be specified and controlled in a cleanroom. Temperature control is necessary for operator comfort and process control. This control is important. A rule of thumb is that chemical reactions, such as an etch rate, change by a factor of 2 with a 10°C change in temperature. A typical temperature range is 72°F, $\pm 2^\circ\text{F}$ (22.2°C, $\pm 1.1^\circ\text{C}$).

The relative humidity is also a critical process parameter, especially in patterning areas. In this area, thin films of a polymer (photoresist) are put on the wafer to act as an etch stencil. If the humidity is too high, the wafers collect moisture, preventing the polymer from sticking. The situation is the same as applying paint on a wet surface. On the other side, low humidity can foster the buildup of static charge on the wafer surface. The charge causes the wafer to attract particles out of the air. Relative humidity is controlled at between 15 and 50 percent.

Smog is another airborne contaminant in a cleanroom. Again the problem is most critical in the patterning areas. A step in the patterning process is similar to photographic film developing, which is a chemical process. Ozone, a major component of smog, interferes with the development process and must be controlled. Ozone is filtered out of the air by installing carbon filters in the incoming air ducts.

Cleanroom Construction

Selection of the clean air strategy is the first step in the design of a cleanroom. Every cleanroom is a trade-off between cleanliness and cost. Whatever the final design, every cleanroom is built on basic principles. The primary design is a sealed room that is supplied with clean air, is built with materials that are noncontaminating, and includes systems to prevent accidental contamination from the outside or from operators.

Construction Materials

The inside of a cleanroom is constructed entirely of materials that are nonshedding. This includes wall coverings, process-station materials, and floor coverings. All piping holes are sealed, and even the light fixtures must have solid covers. Additionally, the design should minimize flat surfaces that can collect dust. Stainless-steel materials are favored for process stations and work surfaces.

Cleanroom Elements

Both the design of a cleanroom and its operation must be set up to keep dirt and contamination from getting into the room from the outside. [Figure 5.15](#) shows a layout for a typical fab-processing area. The following nine techniques are used to keep out and control dirt:

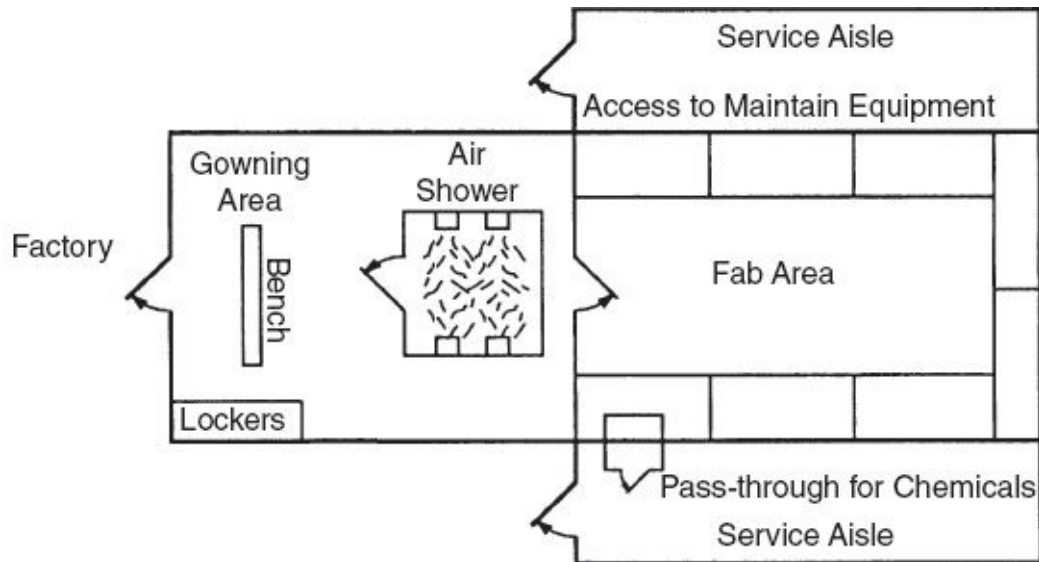


FIGURE 5.15 Fab area with gowning area, air showers, and service aisles.

1. Adhesive floor mats
2. Gowning area
3. Air pressure
4. Air showers
5. Service bays
6. Double-door pass-throughs
7. Static control
8. Shoe cleaners
9. Glove cleaners

Adhesive Floor Mats

At the entrance to every cleanroom is a floor mat with an adhesive surface. The adhesive pulls off and holds dirt adhering to the bottoms of shoes. In some cleanrooms, the entire floor has a surface treated to hold dirt.

Gowning Area

A major part of a cleanroom is the gowning area. This area is a buffer between the cleanroom and the plant. It quite often is supplied with filtered air from ceiling HEPA filters. In this area, the operator's cleanroom apparel is stored. It is also the area where the cleanroom personnel change into their cleanroom garments. The management of this area varies with the degree of cleanliness required in the cleanroom. Quite often, it is managed to the same stringent requirements of the cleanroom itself. Often, the gowning area is divided into two sections by a bench. The operators don the garments on one side and put on their shoe coverings on the bench. The purpose is to keep the area between the bench and the cleanroom at a higher cleanliness level.

The doors between the factory and the cleanroom are never opened at the same time. This procedure ensures that the cleanroom is never exposed directly to the dirtier factory areas. Cleanroom management also includes lists of materials and garments that can and cannot come into the gowning room. Some areas will provide hallway lockers for coats and so on.

Air Pressure

A key design element is the air pressure balance between the cleanroom, the gowning room, and the factory. The well-designed facility will have the air of these three sections balanced such that the highest pressure is in the cleanroom, the second highest in the gowning area, and the lowest in the factory hallways. The higher pressure in the cleanroom causes a low flow of air out of the doors, which, in turn, blows airborne particles back into the dirtier hallway. In self-contained process tools, the loading sections will have positive pressure relative to the general fab area.

Air Showers

The final design element protecting the cleanroom from outside contamination is the air shower located between the gowning room and the cleanroom. Cleanroom personnel enter the air shower, where high-velocity air jets blow off particles from the outside of the garments. An air shower will have an interlocking system to prevent both doors from being opened at the same time.

Service Bays

A cleanroom is really a series of rooms ([Fig. 5.16](#)) within the factory, each contributing to the maintenance of the cleanroom. In the center is the processing cleanroom. Surrounding it is a bay area that is maintained at some designated class number that is generally higher than the cleanroom. Class 1000 or Class 10,000 are typical for service bays. In the bay are the process chemical pipes, electrical power lines, and cleanroom materials. Critical process machines are backed up to the wall dividing the cleanroom and the bay. This arrangement allows technicians to service the equipment from the back without entering the cleanroom.

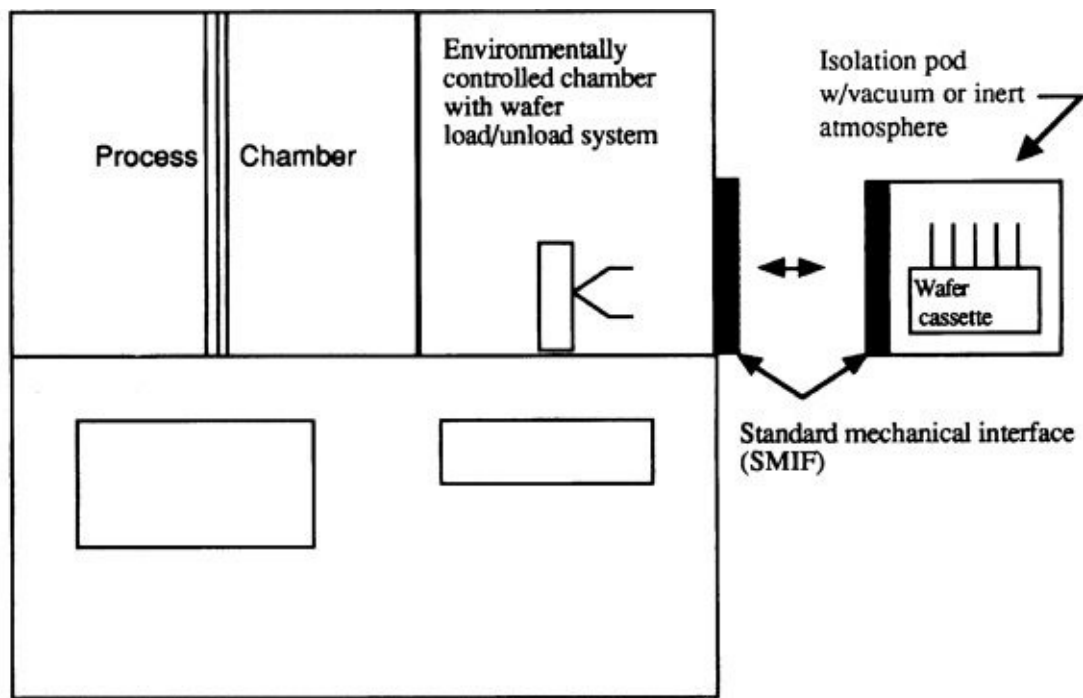


FIGURE 5.16 Minienvironment system elements.

Double-Door Pass-Throughs

The bay also serves as a semiclean area for the storage of materials and supplies. They are put into the cleanroom through double-door pass-through units that protect the cleanliness of the cleanroom. Pass-through units may be simple double-door boxes or may have a supply of positive-pressure filtered air with interlocking devices to prevent both doors from being opened at the same time. Often, the pass-throughs are fitted with HEPA filters. All materials and equipment brought into the cleanroom should be cleaned prior to entry.

Static Control

Higher-density circuits with submicron feature sizes are vulnerable to smaller particles of contamination attracted by static to the wafer. Static charges build up on the wafers, the storage boxes, work surfaces, and equipment. Each of these items can carry static charges as high as 50,000 V (volts) that attract aerosols out of the air and from personnel garments. The attracted particles end up contaminating the wafers. Statically held particulates are very difficult to remove with standard brush and wet-cleaning techniques.

Most static charge is produced by *triboelectric* charging. This occurs when two materials initially in contact are separated. One surface becomes positively charged as it loses electrons. The other becomes negatively charged as it gains electrons. The triboelectric series table in [Fig. 5.17](#) shows the charging potential for some materials found in a cleanroom.¹¹

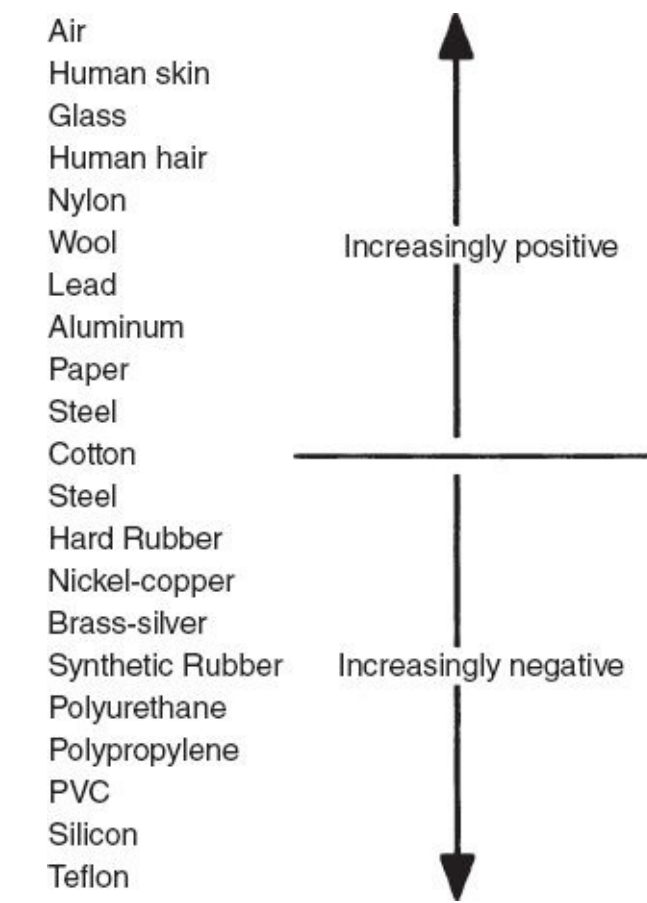


FIGURE 5.17 Triboelectric series. (Source: Hybrid Circuit Technology Handbook, Noyes Publications.)

Static also represents a device operational problem. It occurs in devices with thin

dielectric layers, as in MOS gate regions. An electrostatic discharge (ESD) of up to 10 A (amperes) is possible. This level of ESD can physically destroy an MOS device or circuit. ESD is a particular worry in device-packaging areas. This problem requires that sensitive devices, such as large-array memories, be handled and shipped in holders of antistatic materials.

Photomasks and reticles are particularly sensitive to ESD. A discharge can vaporize and destroy the chrome pattern. Some equipment problems are static related—especially robots, wafer handlers, and measuring equipment. Wafers usually come to the equipment in carriers made of PFA-type materials. This carrier material is chosen for its chemical resistance, but it is not conductive. Charge builds up on the wafers, but cannot dissipate to the carriers. When the carrier comes close to a piece of metal on the equipment, the wafer charge discharges to the equipment. The electromagnetic interference produced interferes with the machine operation.

Static is controlled by prevention of charge buildup and the use of discharge techniques ([Fig. 5.18](#)). Prevention techniques include use of antistatic materials in garments and in-process storage boxes. In some areas, a topical antistatic solution may be applied to the walls to prevent the buildup of static charge. These solutions work by leaving a neutralizing residue on the surface. Generally, they are not used in critical stations because of the possible contaminating effect of the residue.

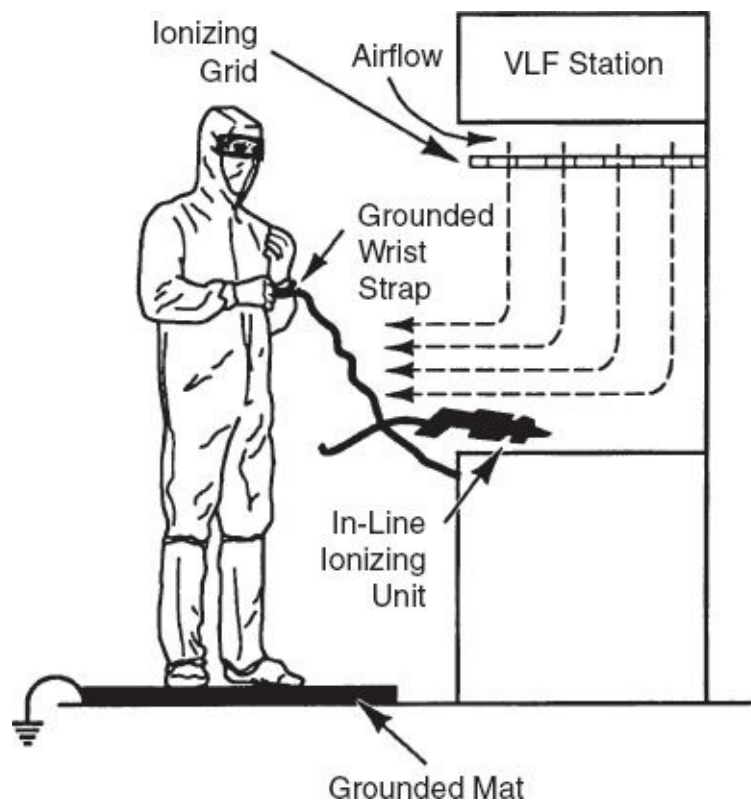


FIGURE 5.18 Static-charge reduction techniques.

Discharge techniques include the use of ionizers and grounded static-discharge straps. Ionizers are placed just underneath the HEPA filters, where they function to neutralize any charge buildup in the filtered air. Ionizers are also placed on nitrogen blow-off guns for the same effect. Some stations will have a portable ionizer blowing ionized air directly on the wafers being processed. Static discharge is also accomplished by grounding operators with wrist straps, having grounded mats at critical stations, and grounding work surfaces. Comprehensive static control programs, with prevention and discharge methods along with personnel training and third-party monitoring, are features in most advanced fabs.

Shoe Covers

In any contamination-controlled area, the dirtiest region is the floor.^{[12](#)} Shoe covers or shoes that never leave the gowning area minimize this contamination source.

Glove Cleaners

Maintaining clean gloves in the fab area is a challenge. Instructing operators to discard gloves when they are contaminated or dirty is one way. However, some contamination is hard to see, and the decision to discard becomes a judgment that can vary from operator to operator. Another approach is to discard the gloves after every shift. This can get very expensive. Some fab areas use glove cleaners that clean and dry the gloves in an enclosure.¹³

Personnel-Generated Contamination

Cleanroom personnel are among the biggest sources of contamination. A cleanroom operator, even after showering and sitting, can give off between 100,000 and 1,000,000 particles per minute.¹⁴ This number increases dramatically when a person is in motion. At 2 mi/hr, a human being gives off up to 5 million particles per minute. The particles come from flakes of dead hair and normal skin flaking. Additional particle sources are hair sprays, cosmetics, facial hair, and exposed clothing. The table in [Fig. 5.19](#) lists the percentage increase in particulate generation over the background contamination level for different operator activities.¹⁵

Normal Breath	No. Particles
Smoker's Breath after Smoking	500%
Sneezing	2000%
Sitting Quietly	20%
Rubbing Hands on Face	200%
Walking	200%
Stamping Foot on Floor	5000%

FIGURE 5.19 Activity-caused increase in particles. (Source: Hybrid Circuit Technology Handbook, Noyes Publications.)

Normal clothing can add millions of particles more to the area, even under a cleanroom garment. In cleanrooms with very high cleanliness levels, operators will be directed to wear street clothing that is made of tight-weave, nonshedding materials. Garments made of wools and cottons are to be avoided, as are ones with high collars.

A human's breath also contains high levels of contaminants. Every exhale puts numerous water droplets and particles into the air. The breath of smokers carries millions of particles for a long time after a cigarette is finished. Body fluids such as

saliva contain sodium, a killer to many semiconductor devices. While healthy human beings are sources of many contaminants, sick individuals are even worse. Specifically, skin rashes and respiratory infections are additional sources of contaminants. Some fabrication areas reassign personnel with certain health problems.

Given the scope of the problem, the only feasible way to render humans acceptable in a cleanroom is to cover them up. The style and material of clothing selected for cleanroom personnel depends on the level of cleanliness required. For a typical area, the clothing material will be nonshedding and may contain conductive fibers to draw off static charges. The trade-off is the filtering ability of the material and operator comfort. The reclean versus discard issue applies to cleanroom suits as well as gloves. Most ultra-large-scale integration (ULSI) fabs have found that reusable gowns, even with the cost of recleaning, provide an overall lowered cost of ownership.¹⁶

Every part of the body is covered up. The head will have an inner cap that keeps the hair in place. This is covered by an outer shell that is designed to fit close to the face and has snaps or a tail for securing the headgear under the body-covering smock. Covering the face will be a mask. Masks vary from surgical types to full ski-mask-style designs. In some cleanrooms, both inner and outer face masks are required. The eyes, which are a major source of fluid particles, are covered by glasses (usually safety glasses) with side shields. In contamination-critical areas, the operators might wear a covering that totally encloses the head and face, connected to an air supply and filter. These are lower-tech models of a space-suit helmet. The unit attaches to a belt filter, blower, and pump system. Fresh air is supplied by the pump, while the filter ensures that no breath-generated contaminants are discharged into the room.¹⁷

Body covers are oversuits (called *bunny suits*) that have closures for the legs, arms, and neck. Well-designed suits will have covers over the zippers and no outside pockets.

The feet are covered with shoe coverings, some with attached leggings that come up the leg. In static-sensitive areas, straps are available to drain off static charge.

Hands are covered with at least one pair of gloves. Most favored are medical-type PVC gloves that permit good tactile feeling. Glove materials for chemical handling include orange latex (acid protection), green nitrile (solvent protection), and silver multilayered PVA solvent (special solvents).¹⁸ In some areas, an inner pair of cotton gloves is permitted for comfort. Gloves should be pulled up over the sleeves to prevent contamination from traveling down the arm and into the cleanroom.

Skin flaking can be further controlled with the use of special lotions that moisten the skin. Any lotions used must be sodium and chlorine-free.

In general, the order of gowning is from the head down. The theory is that dirt

stirred up at each level is covered up by the next lower garment. Gloves are put on last. The garments and procedures needed to control contamination from cleanroom workers are well known. However, the primary level of defense is the dedication and training of the operators. It is easy for an area to become lax in maintaining cleanroom discipline and to suffer high levels of contamination.

Process Water

During the course of fabrication processing, a wafer will be chemically etched and cleaned many times. Each of the etching or cleaning steps is followed by a water rinse. Throughout the entire process, the wafer may spend a total of several hours in water-rinse systems. A modern wafer-fabrication facility may use up to several million gallons of water per day, representing a substantial investment in water processing, delivery to the process areas, and treatment and discharge of wastewater.¹⁹ Given the vulnerability of semiconductor devices to contamination, it is imperative that all process water be treated to meet very specific cleanliness requirements.

Water from a city system contains unacceptable amounts of the following contaminants:

1. Dissolved minerals
2. Particulates
3. Bacteria
4. Organics
5. Dissolved oxygen
6. Silica

The dissolved minerals come from salts in normal water. In the water, the salts separate into ions. For example, salt (NaCl) breaks up into Na^+ and Cl^- ions. Each is a contaminant in semiconductor devices and circuits. They are removed from the water by reverse osmosis (RO) and ion-exchange systems.

The process of removing the electrically active ions changes the water from a conductive medium to a resistive one. This fact is used to improve the quality of deionized (DI) water. Deionized water has a resistivity of $18,000,000 \text{ } \Omega\text{-cm}$ at 25°C . Ultrapure water (UPW) is processed to $18.2 \text{ meg-}\Omega$. [Figure 5.20](#) shows the effect on the resistivity of water when various amounts of dissolved minerals are present.

Resistivity Ohms-cm 25°C	Dissolved Solids (ppm)
18,000,000	0.0277
15,000,000	0.0333
10,000,000	0.0500
1,000,000	0.500
100,000	5.00
10,000	50.00

FIGURE 5.20 Resistivity of water versus concentration of dissolved solids.

The resistivity of all process water is monitored at many points in the fabrication area. The goal and specification is 18 MW in VLSI areas, although some fabrication areas will run with 15-MW water levels. Solid particles (particulates) are removed from the water by sand filtration, earth filtration, and/or membranes to submicron levels. Bacteria and fungi find water a favorable host. They are removed by sterilizers that use ultraviolet radiation to kill the bacteria and filters to remove them from the stream of water.

Organic contaminants (plant and fecal materials) are removed with carbon bed filtration. Dissolved oxygen and carbon dioxide are removed with forced draft decarbonators and vacuum degasifiers.²⁰ Water specifications for a 4-MB DRAM product are shown in [Fig. 5.21](#).

Water quality parameter	Units	Typical municipal water supply	Typical ultrapure water product
Resistivity	M Ohms-cm	0.004	>18
pH	Units	i	6
TOC	ppb	3500	<10
Amonium	ppb	300	<1
Calcium	ppb	22000	<1
Magnesium	ppb	4000	<1
Potassium	ppb	4500	<1
Silica	ppb	4780	<10
Sodium	ppb	29000	<1
Chloride	ppb	15000	<1
Fluoride	ppb	740	<1
Sulfate	ppb	42000	<1

FIGURE 5.21 Ultrapure water specifications.

The cost of cleaning process water to acceptable levels is a major operating expense of a fabrication area. In most fabrications, the process stations are fitted with water meters that monitor the used water. If the water falls within a certain range, it is recycled in the water system for cleanup. Excessively dirty water is treated as dictated by regulations and discharged from the plant. A typical fabrication system is shown in [Fig. 5.22](#). Water stored in the system is blanketed with nitrogen to prevent the absorption of carbon dioxide. Carbon dioxide in the water interferes with resistivity measurements, causing false readings.

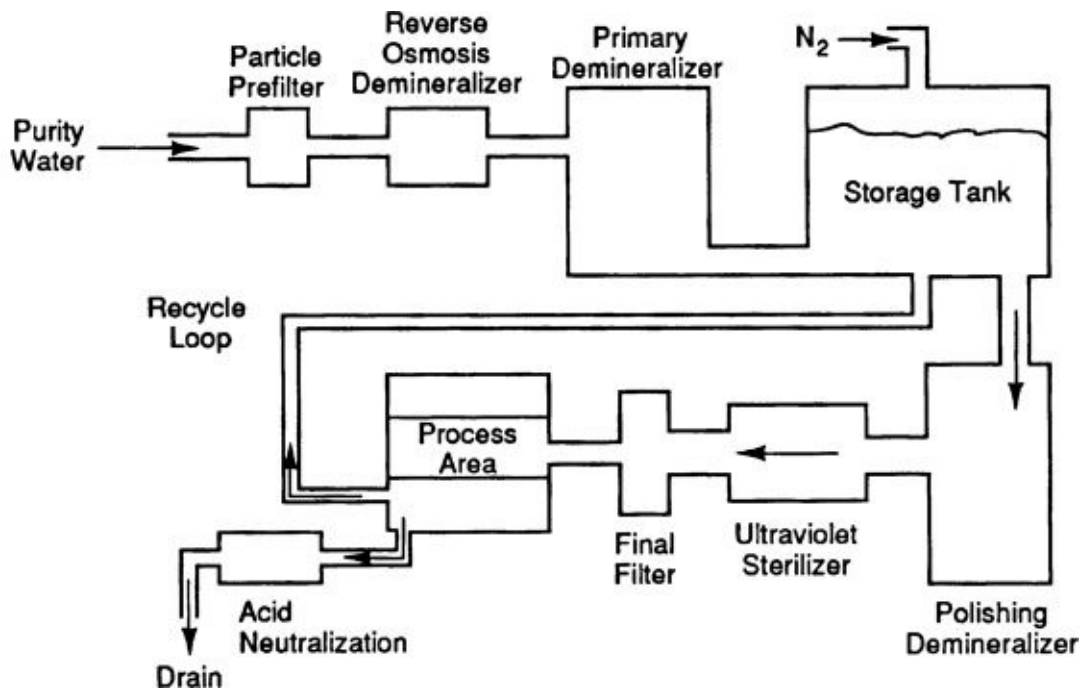


FIGURE 5.22 Typical deionized water system.

Process Chemicals

The acids, bases, and solvents used to etch and clean wafers and equipment have to be of the highest purity for use in a fabrication area. The contaminants of concern are trace metals, particulates, and other chemicals. Unlike water, process chemicals are purchased and used as they come into the plant. Industrial chemicals are rated by grade. They are commercial, reagent, electronic, and semiconductor grade. The first two are generally too dirty for semiconductor use. Electronic grade and semiconductor grade are cleaner chemicals, but levels of cleanliness vary from manufacturer to manufacturer.

Trade organizations such as Semiconductor Equipment and Materials International (SEMI) are establishing cleanliness specifications for the industry. Of primary concern are the metallic mobile ionic contaminants. These are usually limited to levels of 1 part per million (ppm) or less. Some suppliers are making available chemicals with MIC levels of only 1 part per billion (ppb). Particulate filtering levels are specified at 0.2 μm or lower.

Chemical purity is indicated by its assay number. The assay number indicates the percentage of the chemical in the container. For example, an assay of 99.9 percent on a bottle of sulfuric acid means that the bottle contains 99.9 percent sulfuric acid and 0.01 percent other substances.

The delivery of a clean chemical to the processing area involves more than just making a clean chemical. Care must be taken to clean the inside of the containers, using containers that do not dissolve, using particulate-free labels, and placing the clean bottles in bags before shipping. Chemical bottle use is found only in older and lower technology fabs. Most firms purchase clean process chemicals in bulk quantities. They are decanted into smaller vessels and are distributed to the process stations from a central location by pipes or distributed directly to the process station from a mini-unit. Bulk chemical distribution systems (BCDSs) offer cleaner chemicals and lower cost. Special care must be maintained to ensure that piping and transfer vessels are cleaned regularly to prevent contamination.

The requirement for cleaner chemicals, tighter process control, and lower costs is being met by several techniques. One is point-of-use (POU) chemical mixers (another version of a BCDS). These units, connected to a wet bench or automatic machine, mix chemicals and deliver them directly to the process vessel or unit. Another technique is the use of chemical reproprocessors. These units are located on the drain side of a wet process station. Discharged chemicals are refiltered and (in some cases) recharged for reuse. Critical recirculating etch baths are connected to filters to keep a clean supply of chemicals at the wafer. Another approach is point of use chemical generation (POUCG). Chemicals such as NH_4OH , HF , and H_2O_2 are

made at the process station by combining the proper gases with DI water. Elimination of the contamination associated with chemical packaging and transfers gives this method the possibility of producing chemicals in the parts per trillion (ppt) range.²¹

In addition to the many wet (liquid) chemical processes, a semiconductor wafer is processed with many gases. The gases are both the air-separation gases (oxygen, nitrogen, and hydrogen) and specialty process gases such as arsine and carbon tetrafluoride.

Like the wet chemicals, they have to be delivered clean to the process stations and tools. Gas quality is measured in four categories:

1. Percentage of purity
2. Water vapor content
3. Particulates
4. Metallic ions

Extremely high purity is required for all process gases. Of particular concern are the gases used in oxidation, sputtering, plasma etch, chemical vapor deposition (CVD), reactive ion etch, ion implantation, and diffusion. All of these processes involve chemical reactions that are driven by an energy source. If the gases are contaminated with other gases, the anticipated reaction can be significantly altered or the result on the wafer changed. For example, chlorine contamination in a tank of argon used for sputtering can end up in the sputtered film, with disastrous results for the device. Gas purity is specified by the assay number, with typical values ranging from 99.99 to 99.999999 percent, depending on the gas and the use in the process. Purity is expressed as the number of 9s to the right of the decimal point. The highest purity is referred to as “six 9s pure.”²²

The challenge in fab areas is maintaining the gas purity between the manufacturer and the process station. From the source, the gas passes through a piping system, a gas panel containing valves and flowmeters, and a connection to the tool. Leaks in any part of the system are disastrous. Outside air (especially the oxygen) can enter into a chemical reaction with the process gas, changing its composition and the desired reaction in the process. Contamination of the gas can come from outgassing of the system materials. A typical system features stainless-steel piping and valves along with some polymer components, such as connectors and seals. For ultra-clean systems, the stainless-steel surfaces have electropolished and/or double vacuum melt inside surfaces to reduce outgassing.²³ Another technique is the growing of an iron oxide film, called *oxygen passivation* (OP), to further reduce outgassing. Polymer components are avoided. Gas panel design eliminates dead spaces that become repositories for contamination. Also critical is the use of clean welding processes to eliminate gases absorbed into the piping from the welding gases.

The control of water vapor is also critical. Water vapor is a gas and can enter into unwanted reactions just like other contaminating gases. In fabrication areas, processing silicon wafers with water vapor present is a particular problem. Silicon oxidizes easily wherever free oxygen or water is available. Control of unwanted water vapor is necessary to prevent accidental oxidation of silicon surfaces. Water-vapor limits are 3 to 5 ppm.

The presence of particulates and/or metallics in a gas has the same effect on the processing as in wet chemicals. Consequently, gases are filtered to the 0.2- μ m level and metallic ions are controlled to parts per million or lower.

The air-separation gases are stored on the site in the liquid state. In this state, they are very cold, a situation that freezes some contaminants in the bottom of the tanks. Specialty gases are purchased in high-pressure cylinders. Since many of the specialty gases are toxic or flammable, they are stored in special cabinets outside the plant.

Quartz

Wafers spend a lot of process time in quartz holders such as wafer holders, furnace tubes, and transfer holders. Quartz can be a significant source of contamination, both from outgassing and particulates. Even high-purity quartz contains heavy metals that can outgas into diffusion and oxidation tube gas streams, especially during high-temperature processes. Particulates come from the abrasion of the wafers in the wafer boats and the scrapping of the boats against the furnace tubes. Both electric and flame fused quartz processes are used to produce quartz surfaces that are acceptable for semiconductor use.²⁴

Equipment

Successful contamination control is dependent on knowing the sources of contamination. Most analysis ([Fig. 5.23](#)) identifies process equipment as the largest source of particulates. By the 1990s, equipment-induced particulates rose to the level of 75 to 90 percent of all particle sources.²⁵ This does not mean that the equipment is getting dirtier. Advances in particulate control in air, chemicals, and from personnel has shifted the focus to equipment. Defect generation is part of equipment specifications. Generally, the number of particles added to the wafer per product pass (ppp) through the tool is specified. The term *particles per wafer pass* (PWP) is also used.

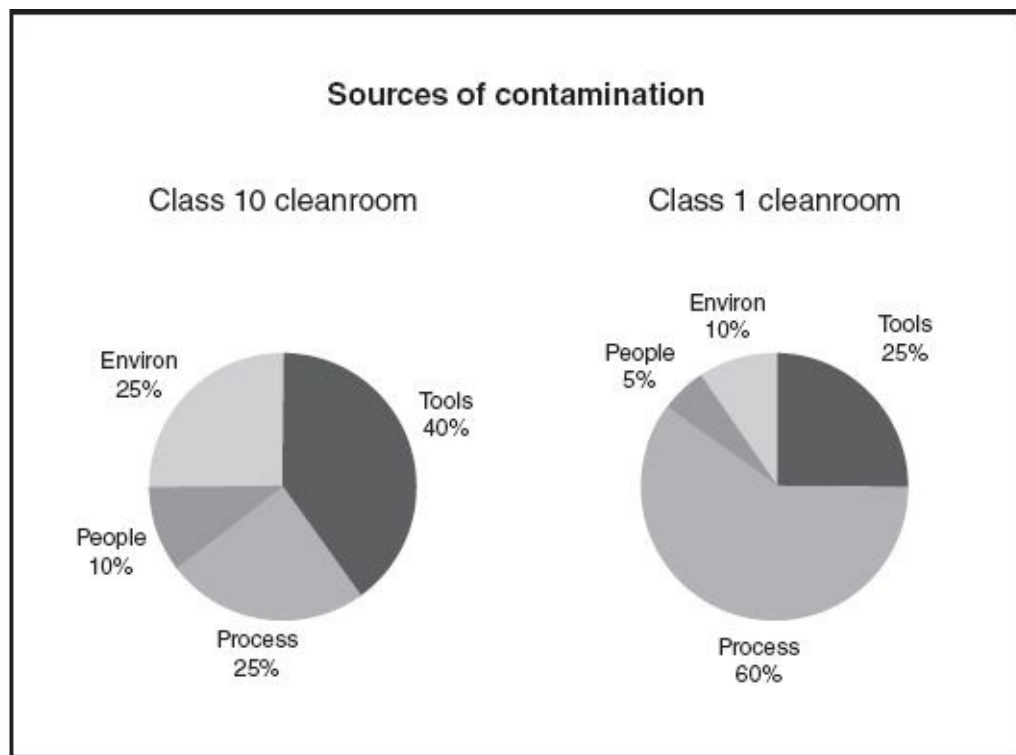


FIGURE 5.23 Sources of particulate contamination. (Courtesy of Mark Jamison, 300-mm Wafer Fab Contamination Control, HDR Architecture, Inc.)

Reduction of particle generation starts with design and material selection. Other factors are the transport of the particles to the wafer and deposition mechanisms such as static. Most equipment is assembled in a cleanroom with the same class number as the customer's fab area. The clustering of several process tools with one clean microenvironment for loading and unloading minimizes the contamination generation associated with multiple loading stations. Cluster tools are discussed in [Chap. 15](#). There is interest in having *in situ* particle monitors in the process

chambers.



Cleanroom Materials and Supplies

In addition to the process chemicals, it takes a host of other materials and supplies to process the wafers. Each of these must meet cleanliness requirements. Logs, forms, and notebooks, if used, will be either of nonshedding coated paper or of a polymer plastic. Pencils are not allowed, and pens are the nonretracting type. Computer monitors are used to maintain logs and process outcomes.

Wafer storage boxes (FOUPs) are made of specific nonparticle-generating materials as are carts and tubing materials. Cart wheels and tools are used without grease or lubricants. In many areas, mechanics' tools and toolboxes are cleaned and left inside the cleanroom.

Cleanroom Maintenance

Regular maintenance of the cleanroom is essential. The cleaning personnel must wear the same garments as the operators. Cleanroom cleaners and applicators, including mops, must be carefully specified. Normal household cleaners are far too dirty for use in cleanrooms. Special care must be taken when using vacuum cleaners. Special cleanroom vacuums with HEPA filtered exhausts are available. Many cleanrooms will have built-in vacuum systems to minimize dirt generation during cleaning.

The wipe-down of process stations is done with special wipes made from nonshedding polyester or nylon, prewashed to reduce contamination. Some are prewetted with an isopropyl alcohol and DI water solution, which provides convenience and eliminates secondary contamination from spraying cleaners in the cleanroom.²⁶

The wiping procedure is critical. Wall surfaces should be wiped from top to bottom and deck surfaces from back to front. Cleaning chemicals in spray bottles, if used, should be sprayed into wipes, not onto the surfaces. This prevents overspray onto the wafers and equipment. Cleanroom cleaning has, itself, become a critical supporting operation. Many fabs use outside certification services²⁷ to identify and document cleanliness levels, practices, documentation, and control procedures. Certification standards for cleanroom maintenance schedules is identified in the ISO Global Cleanroom Standards (ISO 14644-2).

Copper metallization has become the metal of choice for advanced ULSI devices. While the copper has many advantages (see [Chap. 13](#)), there is a serious downside. Copper contamination inside a silicon wafer raises havoc with the electrical operation of the component devices. Isolating the copper-processing areas and copper-deposited wafers is essential. Separate areas and separate process tools are required with strict production control to ensure that no copper technology wafers get diverted into the other process areas.

Wafer-Surface Cleaning

Clean wafers are essential at all stages of the fabrication process, but are especially necessary before any of the operations are performed at high temperature. Up to 30 percent of all process steps relate to wafer cleaning.²⁸ The cleaning techniques described here are used throughout the wafer-fabrication process.

The story of semiconductor process development is in many respects the story of cleaning technology keeping pace with the increasing need for contamination-free wafers. Wafer surfaces have four general types of contamination. Each represents a

different problem on the wafer, and each is removed by different processes. The four types are:

1. Particulates
2. Organic residues
3. Inorganic residues
4. Unwanted oxide layers

In general, a wafer-cleaning process, or series of processes, must (a) remove all surface contaminants (listed above), (b) not etch or damage the wafer surface, (c) be safe and economical in a production setting, and (d) be ecologically acceptable. Cleaning processes are generally designed to accommodate two primary wafer conditions. One, *front end of the line* (FEOL), refers to the wafer-fabrication steps used to form the active electrical components on the wafer surface. During these steps, the wafer surface, and particularly the gate areas of MOS transistors, is exposed and vulnerable. A critical parameter in these cleaning processes is *surface roughness*. Excessive surface roughness can alter device performance and compromise the uniformity of layers deposited on top of the devices. Surface roughness is measured in nanometers as the root mean square of the vertical surface variation (nmRMS). For CMOS devices the variation has dropped to less than 0.1 nm by the year.²⁹ Other concerns in FEOL cleaning processes are the electrical condition of the bare surface. Metal contaminants on the surface change the electrical characteristics of devices, with MOS transistors being particularly vulnerable. Sodium (Na) is a particular problem (see [Chap. 4](#)), along with Fe, Ni, Cu, and Zn. Cleaning processes will have to reduce concentrations to less than 2.5×10^9 atoms/cm² to meet advanced device needs. Aluminum and calcium are also problems and need reduction on the surface to the 5×10^9 atoms/cm² level.³⁰

A most critical factor is maintaining the integrity of gate oxides. Cleaning processes can attack and rough up gate oxides, with the thinner ones being most vulnerable. A gate oxide is required to act as a dielectric in an MOS transistor, and as such must have a consistent structure, makeup, and thickness. *Gate oxide integrity* (GOI) is measured by testing the gate for electrical shorts.³¹

Specific concerns at the BEOL cleans, in addition to particles, metals, and general contamination, are anions, polysilicon gate integrity, contact resistance, via hole cleanliness, organics, and the overall numbers of electrical shorts and opens in the metal system. These issues are explored in [Chap. 13](#). Photoresist removal is also a cleaning process with both FEOL and BEOL consequences. These issues are explored in [Chap. 9](#).

Different chemicals and cleaning methods are mixed and matched to accommodate the needs at particular steps in the process. A typical FEOL cleaning process (such as preoxidation clean) is listed in [Fig. 5.24](#). The FEOL process listed

is called a non-HF-last process. Other variations have the HF removal step last. Non-HF-last surfaces are hydrophilic that can be dried without watermarks and have a thin oxide (grown in the cleaning steps) that can protect the surface. They also absorb more organic contamination. HF-last surfaces are hydrophobic, which can be difficult to dry without watermarks if there are also hydrophilic (oxide) surface layers present. These surfaces are stable from hydrogen surface passivation.³¹ A choice of either non-HF-last or HF-last depends on the sensitivity of the devices being fabricated in the wafer surface and general cleanliness requirements.

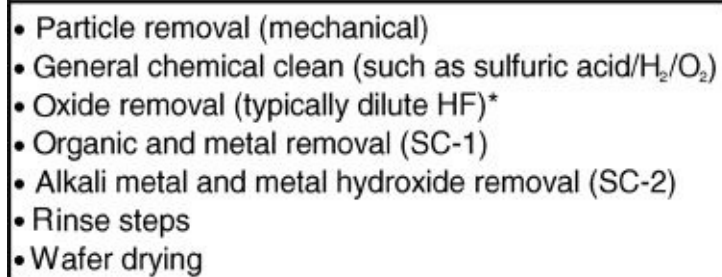
- 
- Particle removal (mechanical)
 - General chemical clean (such as sulfuric acid/H₂/O₂)
 - Oxide removal (typically dilute HF)*
 - Organic and metal removal (SC-1)
 - Alkali metal and metal hydroxide removal (SC-2)
 - Rinse steps
 - Wafer drying

FIGURE 5.24 Typical FEOL cleaning process steps.

Particulate Removal

Particulates on the wafer surface vary from very large ones (50 μm in size) to very tiny ones (<1 μm in size). The larger ones can be cleaned off by conventional chemical baths and rinses. The smaller particulates are more difficult to remove, because they can be held to the surface by several strong forces. One is called the *van der Waals force*. It is a strong interatomic attraction between the electrons of one atom and the nucleus of another. A technique to minimize electrostatic attraction is manipulation of a factor called the *zeta potential*. Zeta potential arises from a charge zone around particles that is balanced by an oppositely charged zone in the cleaning liquid. This charge varies with velocity (the speed of movement of the cleaning liquid, as when the wafers are moved in a cleaning bath), the pH of the solution, and the concentration of electrolytes in the solution. It is also affected by additives to the cleaning solution, such as surfactants. These conditions can be set to create a large charge that has the same polarity of the wafer, thus creating a repulsion that serves to keep the particle in the solution and off the wafer surface.

Capillary force is another problem. It occurs when there is a liquid bridge between a particle and the surface ([Fig. 5.25](#)). Capillary forces can be higher than van der Waals forces.³² Surfactants and mechanical assist, such as megasonics, are used to dislodge these particles from the surface.

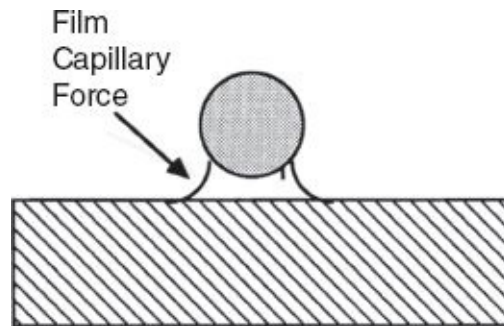


FIGURE 5.25 Capillary force from film.

Cleaning processes are most often a series of steps designed to remove both the large and small particles. The simplest particulate removal process is to blow off the wafer surface using a spray of filtered high-pressure nitrogen from a hand-held gun located in the cleaning stations. In fabrication areas where small particles are a problem, the nitrogen guns are fitted with ionizers that strip static charges from the nitrogen stream and neutralize the wafer surface.

Nitrogen blow-off guns are effective in removing most large particles. Since the guns are hand held, the operators must use them in a manner that does not

contaminate other wafers in the station or the station itself. Blow-off guns are not generally used in Class 1/10 cleanrooms.

Wafer Scrubbers

The stringent wafer cleanliness requirements for epitaxial growth led to the development of mechanical wafer-surface scrubbers, which are used wherever particulate removal is critical ([Fig. 5.26](#)).

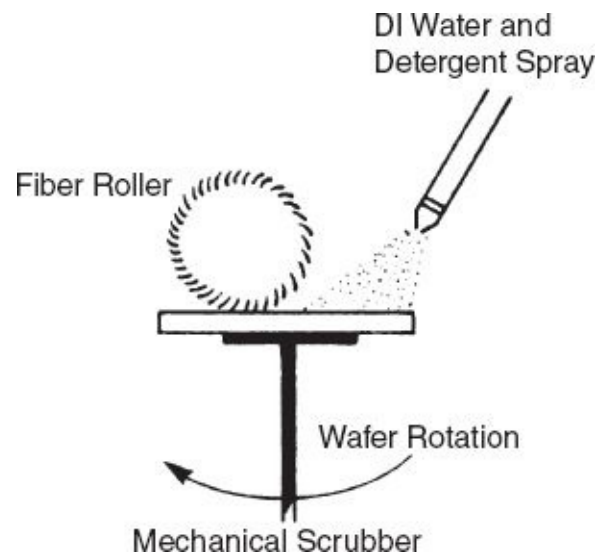


FIGURE 5.26 Mechanical scrubber.

The scrubbers hold the wafer on a rotating vacuum chuck. A rotating brush is brought in near contact with the rotating wafer while a stream of deionized water is directed onto the wafer surface. The combination of the brush and wafer rotations creates a high-energy cleaning action at the wafer surface. The liquid is forced into the small space between the wafer surface and the brush ends where it achieves a high velocity, which aids the cleaning action. Caution must be exercised to keep the brushes and cleaning liquid lines clean to prevent secondary contamination. Also, the brush height above the wafer must be maintained to prevent scratching the wafer surface.

Surfactants may be added to the DI water to increase the cleaning effectiveness and prevent static buildup. In some applications, diluted ammonium hydroxide may be used as the cleaning liquid to prevent buildup of particles on the brush and to control the zeta potential in the system.³³

Scrubbers are designed as stand-alone units with automatic loading capabilities or are built into other pieces of equipment to clean the wafers automatically before processing.

High-Pressure Water Cleaning

The removal of statically attached particles first became a necessity for the cleaning of glass and chrome photomasks. A high-pressure water spray process, uses a small stream of ultrapure water pressurized at 2000 to 4000 psi. The stream is swept across the mask or wafer surface, dislodging both large and small particles. Often, a small amount of surfactant will be added to the water stream to act as a destatic agent.

Organic Residues

Organic residues are compounds that contain carbon, such as oils in fingerprints. These residues can be removed in solvent baths such as acetone, alcohol, or TCE. In general, solvent cleaning of wafers is avoided whenever possible due to the difficulty of drying the solvent completely off the wafer surface. Also, solvents always contain some impurities that themselves may represent a contamination source.

Inorganic Residues

Inorganic residues are those that do not contain carbon. Examples are inorganic acids such as hydrochloric or hydrofluoric acid, which may be introduced from other steps in wafer processing. The organic and inorganic residues are cleaned from the wafers in a variety of cleaning solutions described in the following section.

Chemical-Cleaning Solutions

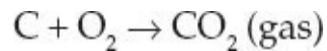
A wide range of cleaning processes exist in the semiconductor industry. Each fabrication area has different cleanliness needs and different experiences with different solutions. The solutions described in this section are those in common use, although there are numerous variations and different combinations of solutions from one fabrication area to another. The processes used to clean wafers before doping, deposition, and metallization steps are described. (The special case of photoresist removal is addressed in [Chap. 8](#).)

Liquid chemical cleaning processes are generally referred to as *wet processes* or *wet cleaning*. Immersion cleaning takes place in glass, quartz, or polypropylene tanks fitted into the deck of a cleaning station (see [Chap. 4](#)). Where heating of the solution is required, the tank may be sitting on a hot plate, be wrapped with heating elements, or have an immersion heater inside. Chemicals are also applied by spraying, either using direct impingement or in centrifugal tools (see spin-rinse dryers in the section “Drying Techniques”).

General Chemical Cleaning

Sulfuric Acid

A common general cleaning solution is hot sulfuric acid with an added oxidant. It is also a general photoresist stripper (see [Chap. 8](#)). The sulfuric acid is an effective cleaner in the 90 to 125°C range. At these temperatures, it will remove most inorganic residues and particulates from the surface. Oxidants are added to the sulfuric acid to remove carbon residues. The chemical reaction converts the carbon to carbon dioxide, which leaves the bath as a vapor by the following reaction:



The oxidant normally used is hydrogen peroxide (H_2O_2) or ozone (O_3). Ozone is also used directly in deionized (DI) water for cleaning and stripping.

Sulfuric Acid with Hydrogen Peroxide

Hydrogen peroxide with sulfuric acid is a common cleaner used to clean wafers at all stages of processing, especially before the tube processes. It is also used as a photoresist stripper in the patterning operation. Within the industry, this solution is known by a number of names, including *Carro's acid* and *piranha etch*. The latter term attests to the aggressiveness and effectiveness of the solution.

A manual method is to add about 30 percent (by volume) of hydrogen peroxide to a beaker of room-temperature sulfuric acid. In this ratio, an exothermic reaction takes place that quickly raises the temperature of the bath into the 110 to 130°C range. As time proceeds, the reaction slows, and the bath temperature falls below the effective range. At this point, the bath may be recharged with additional hydrogen peroxide or discarded. Recharging the bath eventually results in a lowered cleaning rate due to the conversion of the hydrogen peroxide to water, which dilutes the sulfuric acid.

In automated systems, the sulfuric acid is heated to the effective cleaning temperature range and small amounts (50 to 100 mL) of hydrogen peroxide are added before cleaning each batch of wafers. This method maintains the bath at the proper temperature, and the water created from the hydrogen peroxide evaporates out of the solution.

The use of heated sulfuric acid is preferred for economic and process-control reasons. It is also easier to automate this approach to mixing the two chemicals.

Ozone

The purpose of oxidant additives is the addition of oxygen to the solution. Some companies use a flow of gaseous ozone (O_2) directly in the container of sulfuric acid. Ozone and DI water constitute a cleaning solution for light organic contamination. A typical process is 1 to 2 ppm ozone in DI water for 10 minutes at room temperature.³⁴

Oxide Layer Removal

The ease of silicon oxidation has been mentioned. The oxidation can take place in air or in the presence of oxygen in the heated chemical cleaning baths. Often, the oxide grown in the baths, while thin (100 to 200 Å), it is thick enough to block the silicon surface from reacting properly during one of the other process operations. The thin surface oxide can act as an insulator, preventing good electrical contact between the silicon surface and a layer of conducting metal.

The removal of these thin oxides is a requirement in many processes. Silicon surfaces with an oxide are called *hydroscopic*. Surfaces that are oxide-free are termed *hydrophobic*. Hydrofluoric acid (HF) is the acid favored for oxide removal. Prior to an initial oxidation, when the surface is only silicon, the wafers are cleaned in a bath of full-strength HF (49 percent). The HF etches away the oxide but does not etch the silicon.

Later in the processing, when the surface is covered with previously grown oxides, thin oxides in patterned holes are etched away with a water and HF solution. These solutions vary in strength from 100:1 to 10:7 (H₂O to HF). The strength is chosen depending on the amount of oxide already on the wafer, since the water and HF solution will etch both the oxide on the silicon surface in the hole and the oxide covering the rest of the surface. The solution strength is chosen to ensure the removal of the oxide in the holes while not excessively thinning the other oxide layers. Typical dilutions are 1:50 to 1:100.

Management of the chemistry of the silicon surface is an ongoing challenge for cleaning processes. Generally, pregate cleaning for MOS transistors uses the dilute HF as the last chemical step. It is called *HF-last*. HF-last surfaces are hydrophobic and passivated with low metallic contamination. However, hydrophobic surfaces are difficult to dry, often leaving *watermarks*. Another problem is increased particle adhesion and copper plating out of the surface.³⁵

RCA Clean

In the mid-1960s, Werner Kern, an RCA engineer, developed a two-step process to remove organic and inorganic residues from silicon wafers. The process proved to be highly effective, and the formulas became known simply as the *RCA cleans*.³⁶ Whenever an RCA cleaning process is referred to, it means that hydrogen peroxide is used along with some base or acid. The first step, Standard clean 1 (SC-1) uses a solution of water, hydrogen peroxide, and ammonium hydroxide. Solutions vary in composition from 5:1:1 to 7:2:1 and are heated to the 75 to 85°C range. SC-1 removes organic residues and sets up a condition for the desorption of trace metals from the surface. During the process, an oxide film keeps forming and dissolving.

Standard clean-2 (SC-2) uses a solution of water, hydrogen peroxide, and hydrochloric acid mixed in ratios of 6:1:1 to 8:2:1 and is used at the 75 to 85°C temperature range. SC-2 removes alkali ions and hydroxides and complex residual metals. It leaves behind a protective layer of oxide. The original mixtures are shown in [Fig. 5.27](#).

RCA Clean Type	Parts by Volume
Standard Clean (SC-1)	5: DI water 1: 30% Hydrogen peroxide 1: 29% Amonium hydroxide Process 70 degrees centigrade for 5 min.
Standard Clean 2 (SC-2)	6: DI water 1: 30% hydrogen peroxide 1: 37% hydrochloric acid Process 70 degrees centigrade, 5-10 min.

FIGURE 5.27 RCA clean formulas.

The RCA formulas have proven durable over the years and are still the basic cleaning processes for most prefurnace cleaning. Improvements in chemical purity have kept pace with industry cleaning needs. Depending on the application, the order of SC-1 and SC-2 steps may be reversed. Where an oxide-free surface is required, an HF step is used before, in between, or after the RCA cleans.

Many adaptations and changes have been made to the original cleaning solutions. One problem is the removal of metallic ions from the wafer surface. These ions exist in chemicals and are not dissolved (made soluble) in most cleaning and etching solutions. The addition of a *chelating agent*, such as ethylenediamine-tetra-

acetic acid, serves to bind up the ions so they do not redeposit on the wafer.

Diluted RCA solutions are finding more use. SC-1 dilutions are 1:1:50 (instead of 1:1:5) and SC-2 dilutions are 1:1:60 (instead of 1:1:6). These solutions have been found to be as effective as the more concentrated versions. Additionally, they produce less microroughing and are cost-effective and easier to remove.³⁷

Room Temperature and Ozonated Chemistries

The perfect cleaning process takes place at room temperature with totally safe chemicals that are easily and economically disposable. That process does not exist. However, there is research into room-temperature chemistries ([Fig. 5.28](#)). One process³⁸ combines room-temperature baths of ultrapure water infused with ozone and two HF solutions. Megasonic assist elevates the cleaning efficiency.

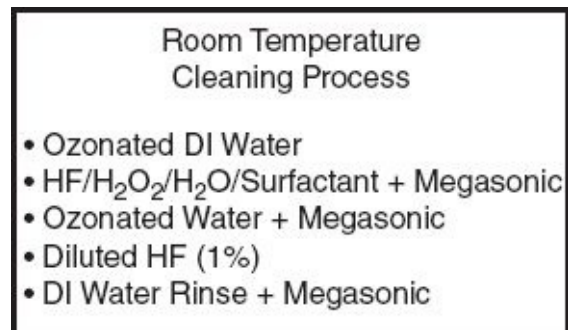


FIGURE 5.28 Experimental room-temperature cleaning process.

Spray Cleaning

The standard cleaning technique has been immersion in chemical baths performed in wet benches or automated machines. While wet cleans are projected to carry the industry into the 0.35-to 0.50- μm era, there are growing considerations. The chemicals are becoming more expensive, immersion in a tank presents a redeposition of contaminants, and smaller and deeper patterns on the wafer surface constrain cleaning efficiency. Alternative cleaning methods have emerged. Spray cleaning offers several advantages. Chemical costs go down, since the chemicals expended are directed to the wafer rather than maintaining large reserves in tanks. Less chemical use reduces the cost of treating or removing hazardous chemical waste. Cleaning efficiency is improved. The pressure of the spray assists in cleaning small patterns on surfaces with deep holes. Also, there is less chance of recontamination, because the wafer receives only fresh chemicals. Spray methods allow immediate water rinsing after cleaning, eliminating transfer to a separate water-rinse station.

Dry Cleaning

The considerations concerning wet immersion methods have spurred interest in and development of *vapor-or gas-phase* cleaning. For cleaning, the wafers are exposed to vapors of the cleaning or etching chemical(s). HF/water vapor mixtures have found their way into processes for oxide removal.³⁹

The ultimate industry dream is all dry cleaning and etching. Dry-etching methods (plasma; see [Chap. 9](#)) are well established. Dry cleaning is under development. Ultraviolet (UV) ozone can oxidize and photo-dissociate contaminants from the wafer surface.

Cryogenic Cleaning

High-pressure carbon dioxide CO_2 , or *snow cleaning*, is a newer technique ([Fig. 5.29](#)). CO_2 is directed at the wafer from a nozzle. As the gas leaves the nozzle, its pressure drops, causing a rapid cooling, which in turn forms either CO_2 particles or snow. The force of the impinging particle dislodges surface particles, and the flow carries them away. The physical bombardment of the surface supplies a cleaning action. Argon aerosol is another cryogenic technique. Argon is a fairly heavy and large atom that can dislodge particles when directed to the wafer under pressure.

A combination process, called *cryokinetic*, combines nitrogen and argon. Precooling of the gases under pressure forms a liquid-gas mixture that is flowed into a vacuum chamber. In the chamber, the liquid expands rapidly to form microscopic crystals that knock particles from the wafer surface.⁴⁰

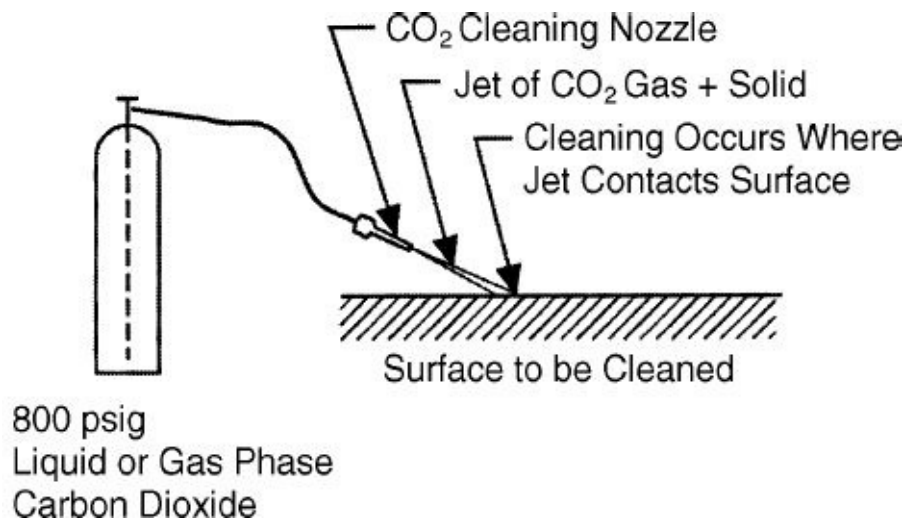


FIGURE 5.29 CO_2 snow cleaning. (Courtesy of Walter Kern.)

Water Rinsing

Every wet-cleaning sequence is followed by a rinse in deionized water. Rinsing performs the dual function of removing the cleaning chemicals from the surface and ending the etching action of an oxide etch. Great quantities of DI water, and great expense, are expended to achieve the goal of clean wafers. Rinsing is done by several different methods. Future focus will be on higher rinse efficiency and an order of magnitude reduction in quantity. The ITRS calls for reducing the present 30 gal/in² of silicon to 2 gal/in² of silicon in the 50-nm gate size sometime around 2012.

Water-rinse techniques include:

- Overflow or cascade rinsers
- Quick dump rinsers (QDRs)
- Ultrasonic or megasonic assist
- Spray
- Spin-rinser dryer tools

Overflow or Cascade Rinsers

The need for atomically clean surfaces precludes just dunking the wafers in a tank of water. Thorough rinsing requires a continuous supply of clean water to the wafer surface. One such method is an overflow rinser ([Fig. 5.30](#)). It is a box sunk into the cleaning station deck. Deionized water is brought into the bottom of the box and flows through and around the wafers, exiting over a dam into a drain system. The rinsing action of the flowing water is enhanced by a stream of nitrogen bubbles introduced into the rinser from the bottom plate. As the nitrogen *bubbles* up through the water, it aids the mixing of the chemicals on the wafer surface with the water. This type of system is often called a *bubbler*. A variation is the *parallel downflow rinser*. In this configuration, the water is brought into the system outside the rinser and directed to flow down through the wafers ([Fig. 5.30](#)).

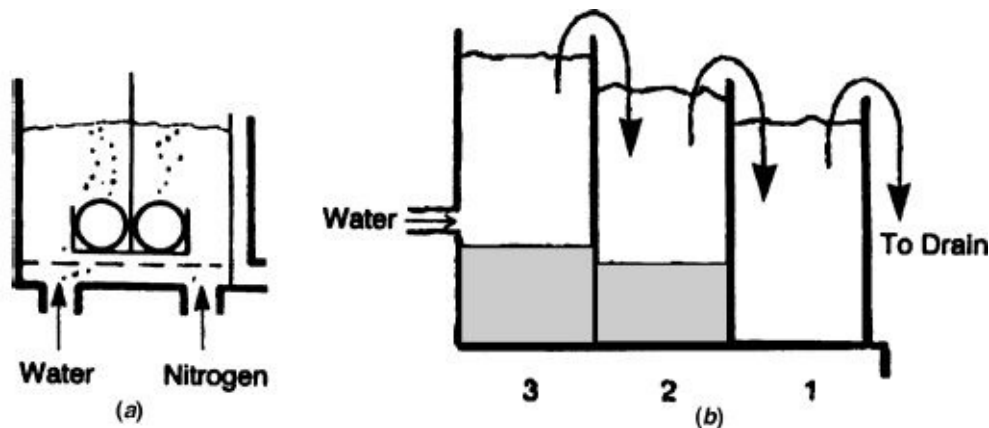


FIGURE 5.30 Rinse systems: (a) Single overflow and (b) three-stage overflow.

A rule of thumb is that adequate rinsing takes a minimum of 5 minutes (depending on wafer diameter) with a flow rate equivalent to five times the volume of the rinser per minute (the number of turnovers per minute). If the size of the bubbler is 3 L, the flow rate should be a minimum of 15 L/min.

The rinse time is determined by measuring the resistivity of the water as it exits the rinser. Cleaning chemicals act as charged molecules in the rinse water, and their presence can be inferred from the electrical resistivity of the water. If ultrapure water (UPW) goes into the rinser at an 18.2-MW level, a reading of 15 to 18.2 MW on the exit side indicates that the wafers are cleaned and rinsed. The rinsing step is so critical that, generally, a minimum of two rinsers are used, and the total rinsing time is set at two to five times the minimum determined by resistivity measurements. Often, a water-resistivity meter is mounted on the outlet to constantly measure the output resistivity and signal when rinsing is completed.

Rinsers

Water-rinse efficiency is both a critical contamination control process and a production cost factor. Parallel downflow rinsing ([Fig. 5.31](#)) has channels that direct the rinse water directly over the wafers, minimizing mixing during the rinse step.

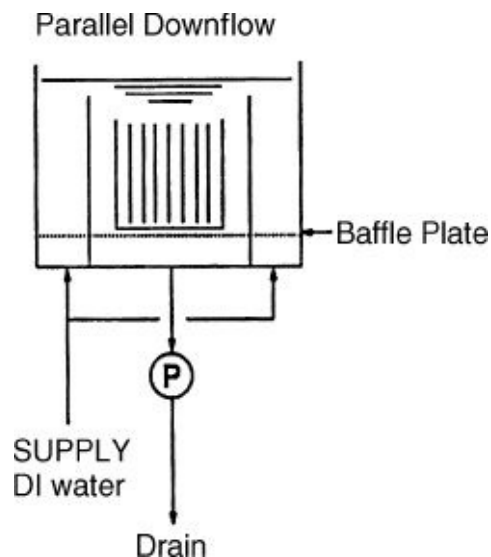


FIGURE 5.31 Parallel downflow rinsing. (Courtesy of Walter Kern.)

With respect to rinsing efficiency and water savings, a dump rinser is also an attractive method. The system is like an overflow rinser but with a spray capability. The wafers are placed into the dry rinser and immediately sprayed with deionized water. While they are being sprayed, the cavity of the rinser is rapidly filled with water. As the water overflows the top, a trapdoor in the bottom swings open, and the water is dumped instantly into the drain system. This fill-and-dump action is repeated several times until the wafers are entirely rinsed.

Sonic-Assisted Cleaning and Rinsing

The addition of sonic energy waves to a tank of cleaning chemicals or a rinse system aids and speeds wet processes ([Fig. 5.32](#)). Use of sonic waves can increase efficiency, which in turn allows a lower bath temperature. Sonic waves are energy waves generated from transducers arranged on the outside of the tank. Two ranges are used. In the 20,000 to 50,000 hertz [1 hertz (Hz) = one cycle per second] range, they are called *ultrasonic* and, in the 850 kilohertz (kHz) range, they are called *megasonic* waves.⁴¹ Ultrasonic assists rinsing by causing a process called cavitation. As the waves pass through the liquid, microscopic bubbles rapidly form and collapse. This creates a microscopic scrubbing action that dislodges particles.

Megasonic assist offers a different mechanism. In fluid flow, there is a static or slow-moving boundary at surfaces, such as wafers. Small particles can be held in this layer, unexposed to the cleaning chemicals. Megasonic energy reduces this layer, exposing the particles to cleaning action. In addition, another phenomenon called *acoustic streaming* fosters an increase in the velocity of the rinse or cleaning solutions passing the wafer surface, increasing cleaning efficiency.⁴²

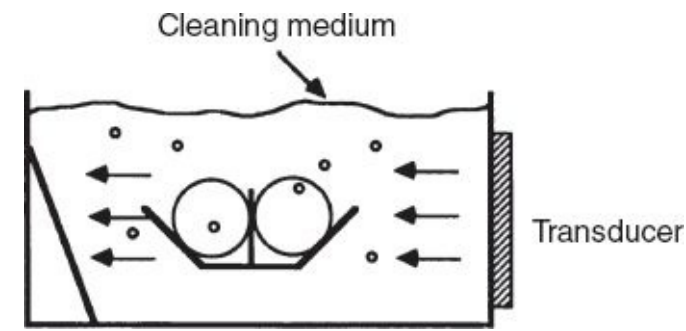


FIGURE 5.32 Ultrasonic/megasonic wafer cleaning/etching bath.

Spray Rinsing

Flowing water removes a water-soluble chemical from the wafer surface by a dilution mechanism. The top layer of the chemical dissolves into the water and is carried away by the flow. This action occurs over and over in a continuous manner. Rinsing is speeded up by a faster water-flow rate, which can remove the dissolved chemical more quickly. The number of turnovers directly determines the rinse rate. This can be understood by considering a very fast water-flow rate, but in a huge rinse tank. The chemical removed from the wafer surface would be evenly distributed throughout the tank, and therefore some would still cling to the surface. The chemical could be removed from the tank only when enough water has flowed in and out of it to carry away the chemical.

One way to have a faster rinse rate is by using a water spray. The spray removes the chemical with a physical force from its own momentum, and the many small droplets continually hitting the wafer surface have the effect of an extremely high turnover rinse. Along with more efficient rinsing, a spray rinser uses considerably less water than an overflow rinser. A problem with a spray rinser comes when the resistivity of the exiting water is measured with a resistivity monitor. Carbon dioxide from the air gets trapped in the water spray; the CO₂ molecules act as charged particles, and the resistivity meter reads them as contaminants, which they are not.

Dump rinsing is also favored, because all of the rinsing takes place in one cavity, which saves equipment and space. It is also a system that can be automated, so the operator only needs to load the wafers in (this can be done automatically) and by the push a button.

Spin-Rinse Dryers

After rinsing, the wafers must be dried. This is not a trivial process. Any water that remains on the surface (even atoms) has the potential of interfering with any subsequent operation. There are three drying techniques (with variations) used ([Fig. 5.33](#)).

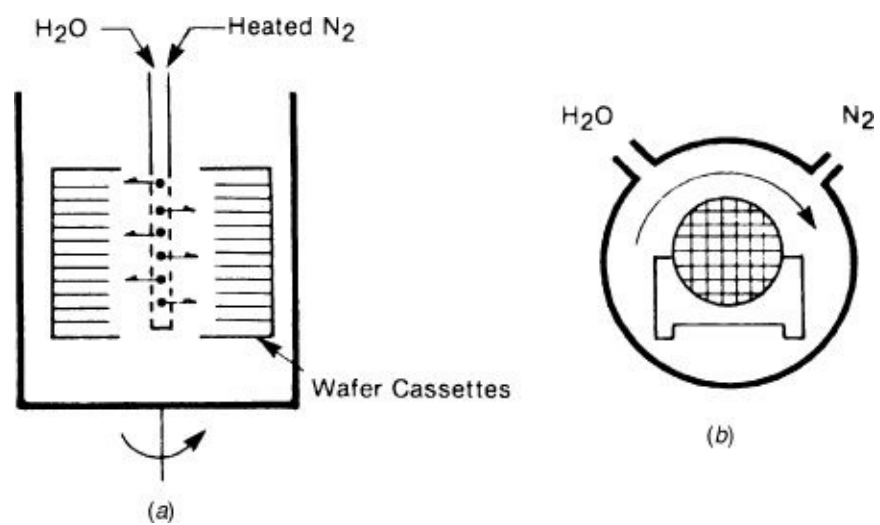


FIGURE 5.33 Spin-rinse dryer styles: (a) multiboat and (b) single-boat axial.

Drying Techniques

- Spin-rinse dryers (SRDs)
- Isopropyl alcohol (IPA) vapor dry
- Surface tension or Marangoni drying

In spin-rinse dryers, complete drying is accomplished in a centrifuge-like piece of equipment. In one version, the wafer boats are put in holders around the inside surface of a drum. In the center of the drum is a pipe with holes, and it is connected to a source of deionized water and hot nitrogen.

The drying process actually starts with a rinse of the wafers as they rotate around the center pipe that sprays the water. Next, the SRD switches to a high-speed rotation as heated nitrogen comes out of the center pipe. The rotation literally throws the water off the wafer surfaces. The heated nitrogen assists in the removal of small droplets of water that may cling to the wafer.

SRDs are also built for drying single-wafer boats. The boat slips into a rotating holder in the center of a chamber. The water and nitrogen come into the chamber through its side rather than through a pipe in the center. The rinsing and drying take place as the boat spins about its own axis. This type of SRD is called an *axial dryer*. These two machines are used for automatic wafer cleaning and etching. As a wafer cleaner, the required chemicals are plumbed to the machine, and microprocessor-controlled valves direct the right chemicals into the chamber.

Isopropyl Alcohol Vapor Drying

A newer drying technique to the semiconductor industry is alcohol drying. In the bottom of the dryer is a heated reserve of liquid IPA with a vapor cloud (vapor zone) above it. When a wafer with residual water on the surface is suspended in the vapor zone, the IPA replaces the water. Chilled coils around the vapor zone condense the water vapor out of the IPA vapors, leaving the wafers water-free. A variation is the *direct displacement vapor dryer*. In this system, the wafers are pulled out of a DI water bath directly into an IPA vapor zone where the water displacement occurs ([Fig. 5.34](#)).

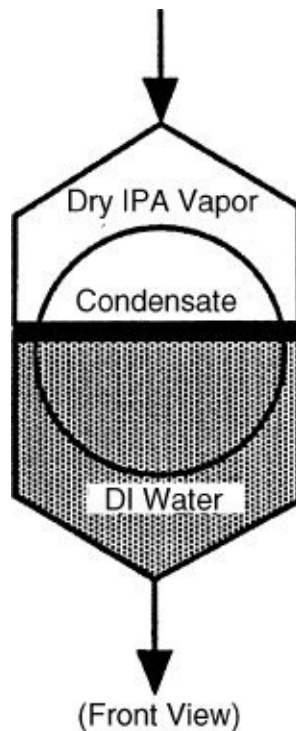


FIGURE 5.34 Vapor dry. (Courtesy of Walter Kern.)

Surface Tension or Marangoni Drying

Water surface tension creates a unique condition when wafers are pulled slowly through a water surface. The tension draws the water away from the surface, leaving it dry. This effect is enhanced when a flow of an organic, such as IPA and nitrogen, is directed at the wafer-water-level interface. The IPA or N₂ flow creates a surface tension gradient, which in turns causes a water flow from the surface into the water. This internal flow further enhances the removal of water from the wafer. In practice, the wafers are either withdrawn from the water bath, or the water is allowed to slowly recede from the rinse tank.^{[43](#)}

Contamination Detection

The detection of the various forms of contamination is detailed in [Chaps. 8](#) and [14](#).

Review Topics

Upon completion of this chapter, you should be able to:

1. Identify the three major effects of contamination on semiconductor devices and processing.
 2. List the major sources of contamination in a fabrication area.
 3. Define the “class number” of a cleanroom.
 4. List the particle density of class 100, 10, and 1 fabrication areas.
 5. Describe the role of positive pressure, air showers, and adhesive mats in maintaining cleanliness levels.
 6. Describe two advantages of contamination control with “wafer isolation.”
 7. List at least three techniques used to minimize contamination from fabrication personnel.
 8. Identify the three contaminants present in “normal” water, and their control in semiconductor plants.
 9. Describe the differences between normal industrial chemicals and semiconductor-grade chemicals.
 10. Name two problems associated with high static levels, and two methods of static control.
 11. Describe a typical FEOL and BEOL wafer-cleaning process.
 12. List typical wafer rinsing techniques.
-

References

- [1.](#) *International Technology Roadmap for Semiconductors*, 1993.
- [2.](#) *International Technology Roadmap for Semiconductors*, Yield Enhancement, 2011.
- [3.](#) Sherry, J., "Assessing Airborne Molecular Contaminant," *Future Fab 9*, International Issue:135.
- [4.](#) "Clean Room and Work Station Requirements, Federal Standard 209E," 1992, Sec. 1-5, Office of Technical Services, Dept. of Commerce, Washington, DC.
- [5.](#) Ibid.
- [6.](#) Semiconductor Industry Association, *National Technology Roadmap for Semiconductors*, 1997, San Jose, CA.
- [7.](#) "Clean Room and Work Station Requirements, Federal Standard 209E," 1992, Sec. 1-5, Office of Technical Services, Dept. of Commerce, Washington, DC.
- [8.](#) Class-10 Technologies, Inc., *Operator Training Course*, 1983, San Jose, CA:13.
- [9.](#) Bonora, A., "Minienvironments and Their Place in the Fab of the Future," *Solid State Technology*, PennWell Publishing, Sep. 1993.
- [10.](#) Newboe, B., "Minienvironments: Better Cleanrooms for Less," *Semiconductor International*, Mar. 1993:54.
- [11.](#) Licari, J., and Enlow, L., *Hybrid Microcircuit Technology Handbook*, 1988, Noyes Publications, Park Ridge, NJ:281.
- [12.](#) "Dryden Engineering Inc., Product Description," Santa Clara, CA:1995.
- [13.](#) Ibid.
- [14.](#) Licari, J., and Enlow, L., *Hybrid Microcircuit Technology Handbook*, 1988, Noyes Publications, Park Ridge, NJ:280.
- [15.](#) Ibid.
- [16.](#) Iscoff, R., "Cleanroom Apparel: A Question of Tradeoffs," *Semiconductor International*, Cahners Publishing, Mar. 1994:65.
- [17.](#) Ibid.
- [18.](#) Ibid.
- [19.](#) Governal, R., "Ultrapure Water: A Battle Every Step of the Way," *Semiconductor International*, Cahners Publishing, Jul. 1994:177.
- [20.](#) Ibid.
- [21.](#) Peters, L., "Point-of-Use Generation: The Ultimate Solution for

Chemical Purity,” *Semiconductor International*, Cahners Publishing, Jan. 1994:62.

[22.](#) Carr, P., “RTP Characterization Using *In-situ* Gas Analysis,” *Semiconductor International*, Cahners Publishing, Nov. 1993:75.

[23.](#) Kobayashi, H., “How Gas Panels Affect Contamination,” *Semiconductor International*, Cahners Publishing, Sep. 1994:86.

[24.](#) Hill, “Quartzglass Components and Heavy-Metal Contamination,” *Solid State Technology*, PennWell Publishing, Mar. 1994:49.

[25.](#) Busnaina, A., “Solving Process Tool Contamination Problems,” *Semiconductor International*, Cahners Publishing, Sep. 1993:73.

[26.](#) Bellville, L., “Presaturated Wipers Optimize Solvent Use,” *Cleanrooms*, Apr. 2000:30.

[27.](#) Gale, G., Kirkpatrick, B., and Kern, F., “Surface Preparation,” *Handbook of Semiconductor Manufacturing Technology*, 2008, CRC Press, New York, NY:Section 5-1.

[28.](#) Allen, R., O’Brian, S., Loewenstein, L., Bennett, M., and Bohannon, B., “MMST Wafer Cleaning,” *Solid State Technology*, PennWell Publishing, Jan. 1996:61.

[29.](#) The Semiconductor Industry Association, *The National Technology Roadmap for Semiconductors*, 1994:116.

[30.](#) Ibid., p. 113.

[31.](#) Kern, W., “Silicon Wafer Cleaning: A Basic Review,” *6th International SCP Surface Preparation Symposium*, 1999.

[32.](#) Steigerwald, J., Murarka, S., and Gutmann, R., *Chemical Mechanical Planarization of Microelectronic Materials*, 1997, John Wiley & Sons, Hoboken, NJ:298.

[33.](#) Hymes, D., and Malik, I., “Using Double-Sided Scrubbing Systems for Multiple General Fab Applications,” *Micro*, Oct. 1996:55.

[34.](#) Burggraaf, P., “Keeping the ‘RCA’ in Wet Chemistry Cleaning,” *Semiconductor International*, Jun. 1994:86.

[35.](#) Kern, W., “Silicon Wafer Cleaning: A Basic Review,” *6th International SCP Surface Preparation Symposium*, 1999.

[36.](#) Ibid.

[37.](#) Lin, F., “Effects of Dilute Chemistries on Particle and Metal Removal Efficiency and on Gate Oxide Integrity,” *5th International Symposium*, SCP Global, 1998.

[38.](#) Wikol, M., “Application of PTFE Membrane Contactors to the Bubble-Free Infusion of Ozone into Ultra-High Purity Water,” *5th International*

Symposium, SCP Global, 1998.

[39.](#) Allen, R., O'Brian, S., Loewenstein, L., et al., "MMST Wafer Cleaning," *Solid State Technology*, PennWell Publishing, Jan. 1996:62.

[40.](#) Butterbaugh, J., "Enhancing Yield through Argon/Nitrogen Cryokinetic Aerosol Cleaning after Via Processing," *Micro*, Jun. 1999:33.

[41.](#) Wolf, S., and Tauber, R. N., *Silicon Processing for the VLSI Era*, p. 519.

[42.](#) Busnaina, A., and Dai, F., "Megasonic Cleaning," *Semiconductor International*, Aug. 1997:85.

[43.](#) Ibid.

44. Wang, J., Hu, J., and Puri, S., "Critical Drying Technology for Deep Submicron Processes," *Solid State Technology*, Jul. 1998:271.

CHAPTER 6

Productivity and Process Yields

Overview

Wafer fabrication and packaging are incredibly long and complex processes involving hundreds of demanding steps. These steps are never performed perfectly every time, and contamination and material variations combine to cause wafer loss in the process. Additionally, some of the individual chips on the wafers fail to meet customer electrical and performance specifications. In this chapter, the major yield measurement points are identified along with the major process and material factors that affect yield. Typical yields for the different yield points and for different circuits are presented.

Yield Measurement Points

Maintaining and improving process and product *yields* is the lifeblood of semiconductor manufacturing. To a casual observer, it would seem that the industry is fixated on production yields. The observation is indeed correct. The demanding nature of the process and sheer number of processes required to produce a packaged chip result in product loss. These two factors result in a production process that typically ships only 20 to 80 percent of the chips it commits into the wafer-fabrication line.

These yields seem extraordinarily low compared to most manufacturing operations. Yet, when one considers the challenge of producing hundreds of circuits, composed of millions of micron or submicron-size patterns in layers that are equally thin, at very stringent cleanliness levels, all within the confines of a 140-mm² chip from some 39 separate masks, it is a testament to the industry that functioning chips are produced at all.¹

Another factor that helps keep yields depressed is the nonrepairable nature of most production mistakes. While defective automobile parts can be replaced, few such options are available in semiconductor manufacturing. Defective chips or wafers generally cannot be recovered. In some cases, chips that fail performance tests can be downgraded and sold for a less demanding use. Scrapped wafers may find a new life as control or monitor wafers (or see discussions of oxidation in [Chaps. 5 and 7](#)).

Added to these process factors is the volume nature of the business. High capital costs and a higher-than-average percentage of engineering personnel translate to a high-overhead business. This high overhead, coupled with competition that keeps downward pressure on selling prices, requires that most chip producers run a high-volume, high-yield process.

Given all of these factors, the preoccupation with yield is understandable. Most suppliers of equipment and materials tout the yield improvements possible with their products. Likewise, process engineering groups have as their prime responsibility the maintenance and improvement of process yields. Yield measurement starts at the individual process level and is tracked through the entire process sequence, from incoming blank wafer to shipment of the completed circuit.

Typically, a plant will monitor yields at three major points in the process. They are at the completion of the wafer-fabrication processes, after wafer sort, and at the completion of the packaging and final test processes ([Fig. 6.1](#)).

Major yield measurement points	
Manufacturing stage	Measured
Wafer fabrication- Yield =	$\frac{\# \text{ Wafers out}}{\# \text{ Wafers started}}$
Wafer sort- Yield =	$\frac{\# \text{ Functioning (good) die}}{\# \text{ Die on wafer(s)}}$
Packaging- Yield =	$\frac{\# \text{ Packages die passing final test}}{\# \text{ Good die started into packaging}}$

FIGURE 6.1 Major yield measurement points.

Accumulative Wafer-Fabrication Yield

A major yield measurement point is at the completion of wafer fabrication. This yield is called by a variety of names, including fab yield, line yield, cumulative fab yield, or “cum” yield.

Whatever the name, it is expressed as a percentage of the wafers leaving the wafer fab divided by the number that entered the process. Since different product types have different components, feature sizes, and density factors, the wafer-fab yield is calculated for each product type rather than the entire fabrication line yield.

The cum fab yield starts by first calculating the number of wafers that leave each of the individual processes (called *station yields*) and dividing by the number that entered the station.

$$\text{Station yield} = \frac{\text{numbers of wafers leaving station}}{\text{numbers of wafers entering station}}$$

The station yields are, in turn, multiplied together to calculate the overall cum fab yield.

$$\text{Cum fab yield} = Y(\text{station 1}) \times Y(\text{station 2}) \times \dots \times Y(\text{station } n)$$

Figure 6.2 lists an 11-step process such as the one illustrated in Chap. 5. Typical station yields are listed in Col. 3 and the accumulated yield in Col. 5. For a single product, the cum fab yield calculated from the station yields is the same as the yield calculated by dividing the number of wafers *out* of fab by the number of wafers started into the fab line. The accumulated yield equals the simple cum fab yield calculation for this individual circuit. Note that, even with very high individual station yields, the cum fab yield will continue to fall as the wafers come through the

process. A modern integrated circuit will require 300 to 500 individual process steps, which represents a huge challenge to maintain profitable productivity. Successful wafer-fabrication operations must achieve accumulative fabrication yields over 90 percent to stay profitable and competitive.

Step	Wafers in	Yield*	Wafers out	Accumulated yield
1. Field oxidation	1000	99.5	995	99.5
2. Source/drain mask	995	99.0	965	96.5
3. Source/drain doping	965	99.3	978	97.8
4. Gate region mask	978	99.0	968	96.8
5. Gate oxidation	968	99.5	964	96.4
6. Contact hole mask	964	94.0	906	90.6
7. Deposit metal layer	906	99.2	899	89.9
8. Metal layer mask	899	97.5	876	87.6
9. Alloy metal layer	876	100	876	87.6
10. Passivation layer deposition	876	99.5	872	87.2
11. Passivation layer mask	872	98.5	859	85.9
*Yield values are typical for the particular steps.				

FIGURE 6.2 Accumulated (wafer-fab) yield calculation.

Wafer-fabrication cum yields vary from 50 to 95 percent, depending on a number of factors. The calculated cum yield is used for production planning and by engineering and management and as a measure of the process effectiveness.

Wafer-Fabrication Yield Limiters

Wafer-fabrication yield is limited by a number of factors. The five listed below are fundamental factors that must be controlled in any wafer-fabrication facility. These basic factors, in combination with device or circuit-specific factors, result in the overall yield of good chips out of a given facility.

1. Number of process steps
2. Wafer breakage and warping
3. Process variation

4. Process defects

5. Mask defects

Number of Process Steps

In the calculation in [Fig. 6.2](#), note that each individual process operation yield must be in the high 90 percent range to produce the 85.9 percent cum fab yield. Illustrated is a fairly simple 11-step process. Ultra-large-scale integration (ULSI) circuits require hundreds of major process operations. Processes with many hundreds of operations are typical for state of the art products.² Each operation requires several steps, each of which in turn involves a number of substeps. It is easy to appreciate the continual pressures on fabrication areas to maintain high cum yields generated by the number of process steps. The more complicated the circuit, with a high number of steps, the lower the expected cum yield.

More process steps also increase the probability that one of the other four yield limiters will affect the wafer during the process. This factor is a tyranny of numbers. For example, to achieve a 75 percent accumulated fabrication yield with a 50-step process, each of the individual steps would have to be 99.4 percent. A further tyranny of this type of calculation is that the cum fab yield can never exceed the lowest individual step yield. If one process step can achieve only a 50 percent yield, the overall cum yield can never be higher than 50 percent.

For each major process operation, there are a number of steps and substeps. In the illustrated 11-step process, the first operation is oxidation. A simple oxidation process requires several steps: cleaning, oxidation, and evaluation. Each of the steps requires substeps. [Figure 6.3](#) lists the six substeps for a typical oxidation process. Each substep represents an opportunity to contaminate, break, or damage the wafer. Automation and isolation technology has provided more control of the wafer environments, yet each transfer and new process environment provides an opportunity for contamination and defects.

Substep	Number of Wafer Handlings
1. Wafers are removed from carrier and placed in cleaning boat.	2
2. Wafers are cleaned, rinsed, and dried.	1
3. Wafers are removed from cleaning boat, inspected, and placed on oxidation boat.	2
4. Boat is removed from furnace.	0
5. Wafers are removed from boat and placed back in carrier.	1
6. Test wafers are removed from carrier and measured.	2
TOTAL NUMBER OF HANDLINGS	8

FIGURE 6.3 Substeps of the oxidation process.

Wafer Breakage and Warping

During the course of the fabrication process, the wafers are handled many times by a combination of manual and automatic techniques. Each handling presents an opportunity to break the relatively fragile wafers. A typical 300-mm (12-in) diameter wafer is only about 800 μm thick. Careful wafer handling is required, and automatic handlers must be maintained to minimize breakage.

Heat treatments add to the susceptibility of the wafers to breaking. Strains are induced in the crystalline material, which make the wafers vulnerable to breaking in subsequent steps. Automatic processing machines accommodate only full-diameter wafers. Therefore, any breakage, however small, is a cause for rejecting the wafer from the process.

Silicon wafers are relatively easy to handle with good practices, and automatic equipment has reduced wafer breakage to a low level. Gallium arsenide wafers, however, are not that resilient, and breakage is a major wafer-yield limiter. In gallium arsenide fabrication lines, where the circuits command a high-selling price, partial wafers are often processed.

Along with minimizing breakage, the wafer surfaces must remain flat throughout the processing. This is especially true on fabrication lines that use patterning techniques that project the pattern onto the wafer surface. If the surface is warped or wavy, the projected image will become distorted and change the required image dimensions. This is similar to projecting a slide onto a distorted screen. Warping comes about from rapid heating and/or cooling of the wafers in tube furnaces (see [Chap. 7](#)).

Process Variation

As the wafer comes through the fabrication process, it receives a number of doping, layering, and patterning processes, each of which must meet incredibly stringent physical and cleanliness requirements. But even the most sophisticated processes vary from wafer to wafer, batch to batch, and day to day. When a process exceeds its process limits (goes out of spec), it will cause some unallowed result on the wafer or within the chips on the wafer.

A goal of process-engineering and process-control programs is not only to keep each process operating within its control specifications but to maintain a constant distribution of the process parameters, such as time, temperature, pressure, and others. These process parameters are monitored with statistical *process control techniques* explained in [Chap. 15](#).

Throughout the process, there are a number of inspections and tests designed to detect unwanted variations as well as frequent calibration of the equipment parameters to process specifications. Some of these tests are performed by production personnel and some by quality control organizations. However, even the best maintained and monitored process exhibits some variations. One of the challenges of process engineering and circuit design is to accommodate the variations and still have a functioning device.

Process Defects

Process defects are defined as isolated regions (or spots) of contamination or irregularities on the wafer surface. These defects are often called *spot defects* or *point defects*. They occur randomly on the wafer surface. Some are nonfatal, and some will render the circuit inoperable. The latter are called *killer defects* (see [Fig. 6.11](#)). Unfortunately, smaller defects are sometimes not detectable during the fabrication process. They become evident at wafer sort as rejected chips.

The major sources of these defects are the various liquids, gases, room air, personnel, process machines, and water used in the fabrication area. Particulates and other small contaminants become lodged in or on the wafer surface. Many of these defects occur in the patterning process. Recall that the patterning process requires using a thin, fragile layer of photoresist to protect the wafer surface during the etch steps. Any holes or tears in the photoresist layer from particulates will end up as tiny etched holes in the wafer surface layer. These holes are called *pinholes* and are a major concern of photomasking engineers. Consequently, the wafers are inspected often for contamination, usually after each major step for contamination. Wafers that exceed the established allowable density are rejected. The SIA *International Technology Roadmap for Semiconductors* (ITRS) calls for maximum defect densities on 300-mm wafer surfaces of 0.68 per square centimeter (cm^2).

Mask Defects

A photomask or reticle is the source of the pattern that is transferred to the wafer surface in the patterning process. Defects on the mask or reticle end up on the wafer as defects or pattern distortions. There are three common mask-or reticle-originated defects. First is contamination, such as dirt or stains on the clear part of the mask or reticle. In optical lithography, they can block the light and print onto the wafer as though they were an opaque part of the pattern. Second are cracks in the quartz plate of the reticle. They also, can block the patterning light and/or scatter the light, causing unwanted images and/or distorted images. Third are pattern distortions that occur in the mask-or reticle-making process. These include pinholes or chrome spots, pattern extensions or missing parts, breaks in the pattern, or *bridges* between adjacent patterns ([Fig. 6.4](#)). Control of mask-generated defects is more critical for device or circuits with smaller feature sizes, higher densities, and larger die sizes.

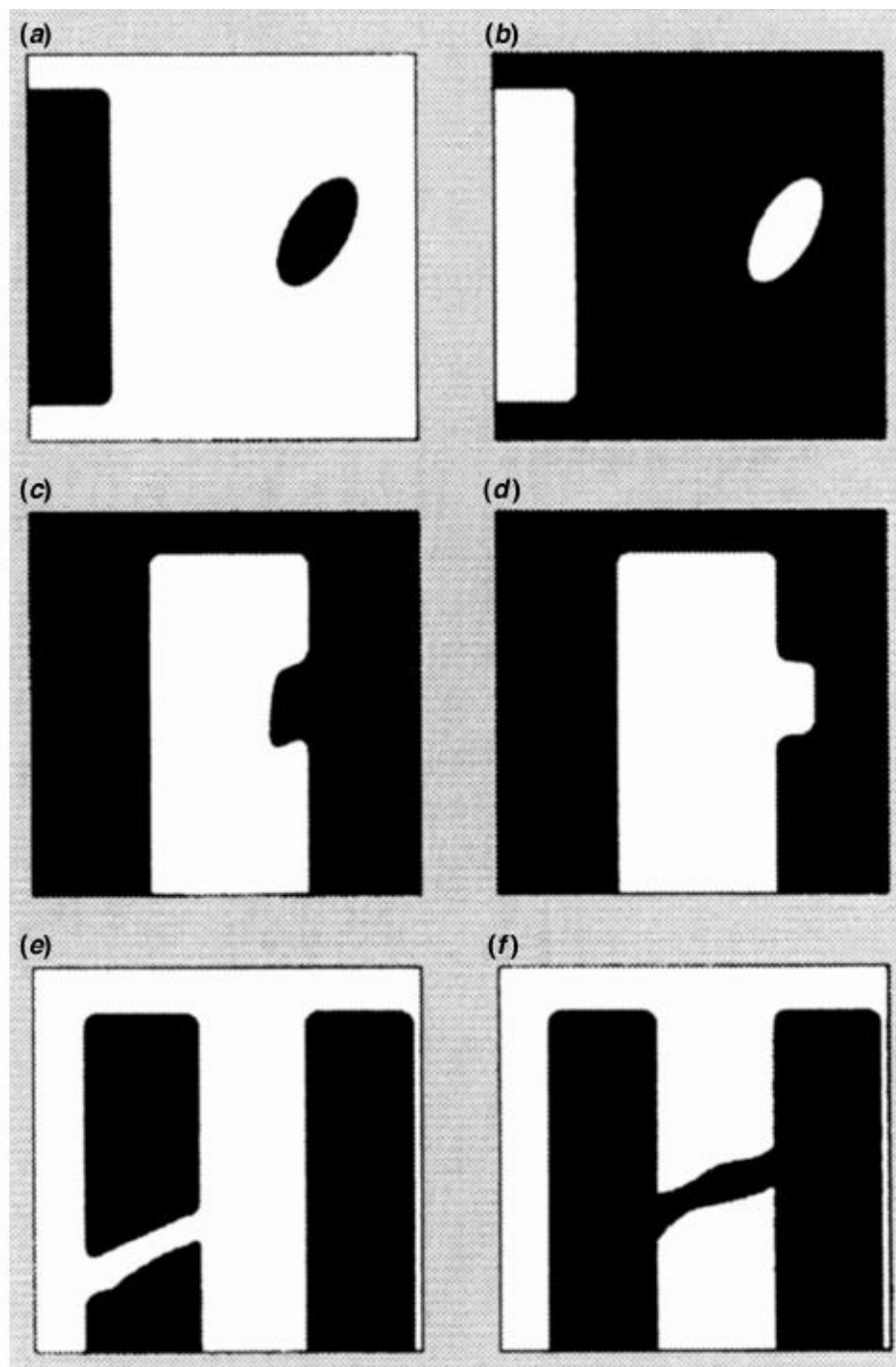


FIGURE 6.4 Mask defects: (a) Spot, (b) hold, (c) inclusion, (d) protrusion, (e) break, and (f) bridge. (Source: Solid State Technology, July 1993, p. 95.)

Wafer-Sort Yield Factors

After fabrication, the wafers go to the wafer sort tester. During the test, each chip will be tested electrically for device specifications and functionality. Up to several hundred individual electrical tests may be performed on each circuit. While these tests measure the electrical performance of the device(s), they indirectly measure the precision and cleanliness of the fabrication processes. Because of natural

process variations and undetected defects, the wafer may have passed all the in-process checks and still have chips that do not function.

Since wafer sort is a comprehensive test, many factors influence the yield. They are: 1. Wafer diameter

2. Die size (area)
3. Number of processing steps
4. Circuit density
5. Defect density
6. Crystal defect density
7. Process cycle time

Wafer Diameter and Edge Die

The semiconductor industry went to round wafers with the introduction of silicon. The first wafers were less than 1 inch in diameter. Since that time, there has been a regular progression to larger-diameter wafers, with 150 mm (6 in) being the standard of VLSI fabrication lines in the late 1980s, and 200-mm wafers being developed for production use in the 1990s. By 2012, 300-mm wafers were in full production with 450-mm diameter wafers being introduced with implementation in full production expected by 2018.³

The move to larger-diameter wafers has been driven by production efficiency, increasing die sizes, and the influence on the wafer-sort yield. Production efficiency is easily understood when one considers that there is an incremental increase in the cost of processing a larger-diameter wafer, while the number of available whole chips on the wafer can increase substantially as illustrated in [Fig. 6.5](#).

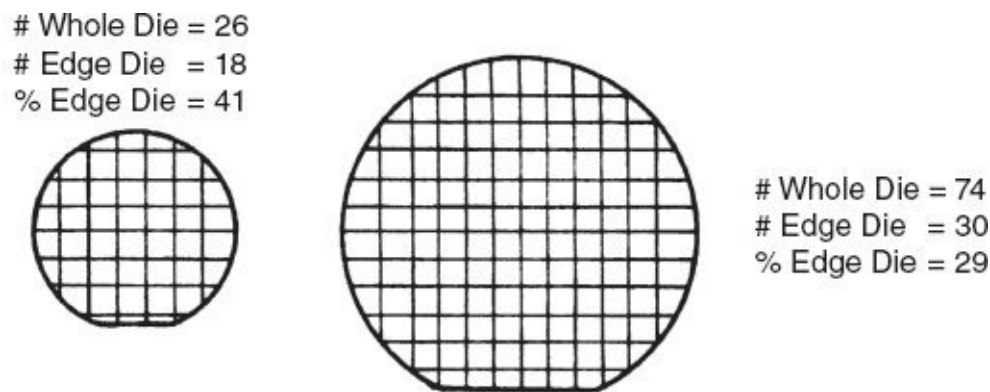


FIGURE 6.5 Effect of larger wafer diameter on percentage of partial die.

The effect of increasing the wafer diameter also has positive effects on the wafer-sort yield. [Figure 6.6](#) shows two wafers of the same diameter but with different die sizes. Note that the smaller-diameter wafer has a very large proportion of its surface covered with partial die—die that cannot function. The larger-diameter wafer, with its greater number and percentage of whole die, if all other factors are equal, will have a higher wafer-sort yield.

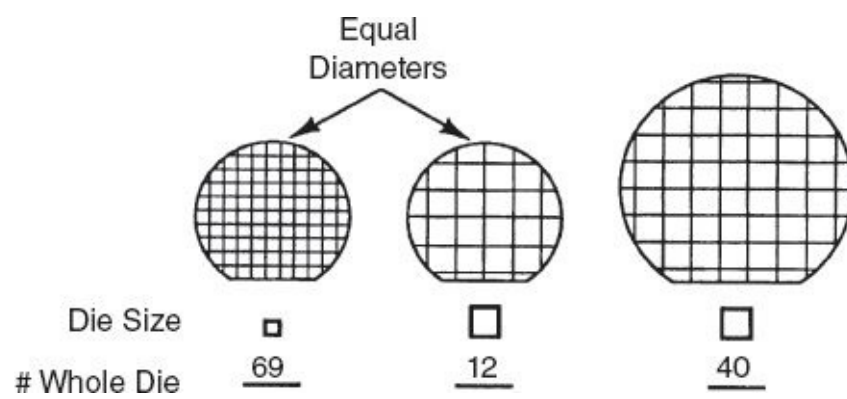


FIGURE 6.6 Whole die count versus die size and wafer diameter.

Wafer Diameter and Die Size

Another driving force for larger-diameter wafers is the trend to larger die sizes. As shown in [Fig. 6.6](#), increasing the die size *without* increasing the wafer diameter also results in a wafer surface with a smaller percentage of whole die. Maintaining a decent wafer-sort yield as the die size increases requires increasing the wafer diameter. [Figure 6.7](#) lists the number of various size chips that will fit on different size wafers. The bottom line is that larger-diameter wafers are more cost-effective.

Die size (mills)	Wafer diameter (mm)			
	150	200	300	450*
300 x 300	293	531	820	1845
400 x 400	113	165	550	1238
500 x 500	108	191	410	923

*Based on 2.25% increase

FIGURE 6.7 Die size versus number of die on wafer.

Wafer Diameter and Crystal Defects

In [Chap. 3](#), the concept of a crystal dislocation was introduced. A crystal dislocation is a point defect *in the wafer* that comes from a local discontinuity of the crystal structure. Dislocations exist throughout the crystal structure and, like contamination and process defect density, affect the wafer-sort yield.

Dislocations also are generated during the fabrication process. They generate (or nucleate) at sites where there are chips and abrasions on the edge of the wafer. These chips and abrasions come from poor handling techniques and automatic handling equipment. The abraded area causes a crystal dislocation. Unfortunately, the dislocation is propagated into the center of the wafer ([Fig. 6.8](#)) during subsequent heat treatments, such as oxidations and diffusions. The length of the dislocation line into the interior of the wafer is a function of the thermal history of the wafer. Consequently, wafers receiving more process steps and/or more heating steps will have more and longer dislocation lines affecting a greater number of chips. One obvious solution to the problem is larger-diameter wafers, which leaves a larger number of unaffected die in the center of the wafer.

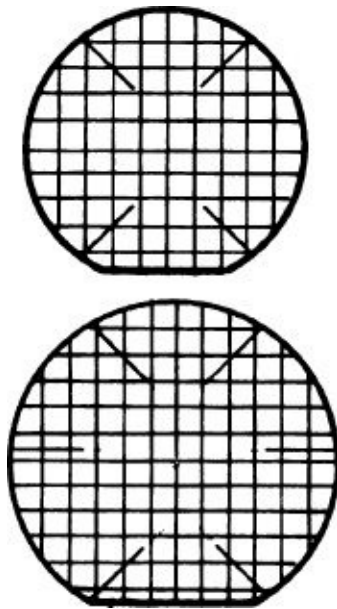


FIGURE 6.8 Effect of dislocations on wafer-sort yield for different wafer dimensions.

Wafer Diameter and Process Variations

The process variations discussed in the section of wafer-fabrication yields affect the wafer-sort yield. In the fabrication area, process variations are detected by sampling inspection and measurement techniques. The nature of inspection sampling is that not all of the variations and defects are detected, so that wafers are passed on with some number of problems. These problems show up at wafer sort as failed devices.

Process variations occur at a higher rate around the edge of the wafer. In the high-temperature processes performed in tube furnaces, there is always some temperature nonuniformity across the wafers. The change in temperature results in uniformity differences on the wafer. Variations occur more at the wafer outer edges where heating and cooling occur at a faster rate. Another contributor to this wafer-edge phenomenon is contamination and physical abuse of the wafer layers that emanates from handling and touching the wafers on their edges. In the patterning process, there can be feature size uniformity problems in the mask-driven processes (full mask projection, proximity and contact exposure). The nature of the light systems is such that the center will be of higher uniformity than the outside edges. In the reticle-driven masking processes (steppers), there is a smaller area of exposure (one or several die), which reduces the image variations across the wafer.

All of these problems result in a lower wafer-sort yield around the edge of the wafer, as illustrated in [Fig. 6.9](#). Larger-diameter wafers help to maintain wafer-sort yields by having a larger area of unaffected die in the center of the wafer.

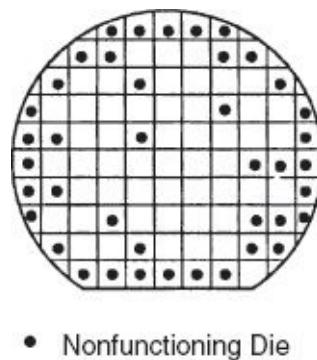


FIGURE 6.9 Typical location of nonfunctioning die after wafer sort.

Die Area and Defect Density

The die size also affects wafer-sort yield relative to the defect density on the wafer surface. The relationship is illustrated in [Fig. 6.10](#). In [Fig. 6.10a](#), a wafer is shown with five defects and no die pattern. This situation illustrates a background defect density produced by all of the fab area factors, regardless of the die size, device type, process control requirements, and so forth. The wafers in [Fig. 6.10b](#) and [6.10c](#) illustrate the effect of this background defect density on the wafer-sort yield for two different die sizes. The larger the die size for a given defect density, the lower the yield.

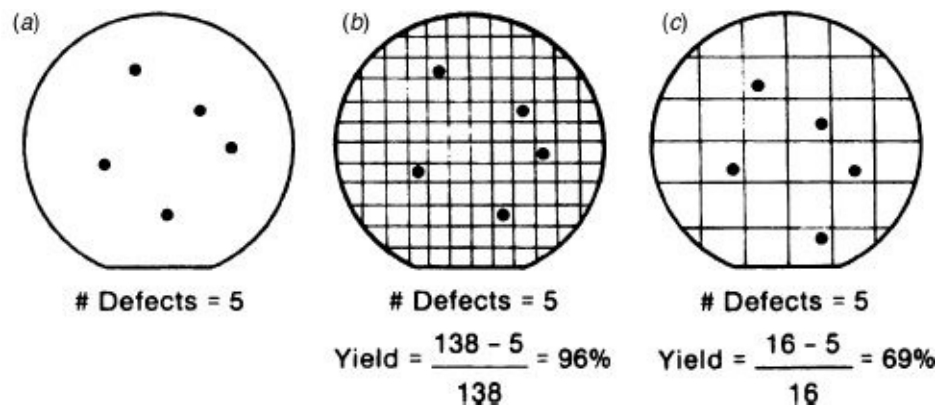


FIGURE 6.10 Effect of background defects on wafer-sort yield for different die sizes.

Circuit Density and Defect Density

The defects on the wafer surface result in die failures by causing a malfunction of some part of the die. Some of the defects are located in nonsensitive parts of the die and do not cause a failure. However, the trend is toward higher levels of circuit integration, which came about because of smaller feature size and a higher density of die components. The result of these trends is a higher probability that any given defect will be in an active part of the circuit, thus lowering the wafer-sort yield as illustrated in [Fig. 6.11](#).

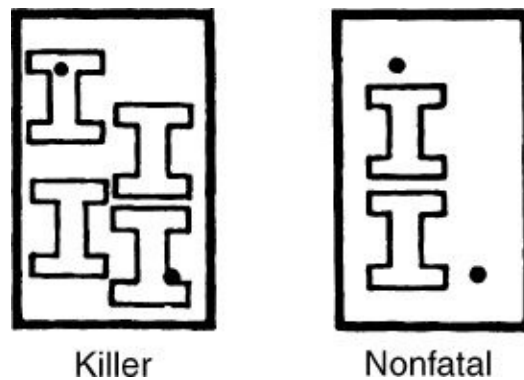


FIGURE 6.11 Killer defects (failed die) and nonfatal (passed die) defects.

Number of Process Steps

The number of process steps was indicated as a limiter of the fab cum yield. The more steps, the greater the opportunities to break or misprocess a wafer. The effect also influences the wafer-sort yield. As the number of process steps increases, the background defect density increases, unless procedures are implemented to lower it. A higher background defect density affects more chips, lowering the wafer-sort yield.

Feature Size and Defect Size

Smaller feature sizes make maintaining an acceptable sort yield difficult from two major factors. First, the smaller images are more difficult to print (see the “Mask Defects” section and [Chap. 8](#)). Second, the smaller images are vulnerable to ever-smaller defect sizes as well as the overall defect density. The 10:1 rule of minimum feature size to allowable defect size has been discussed. One assessment is that, at a defect density of one defect per cm^2 , a circuit with 0.35- μm feature size will have a wafer-sort yield 10 percent less than that of a 0.5- μm circuit processed under the same conditions.⁴

Process Cycle Time

The time that the wafers are actually being processed can be measured in days. But due to queuing at the process stations and temporary slowdowns due to process problems, the wafer often stays in the fab area for several weeks. The longer the wafer is sitting around, the more opportunity for contamination that lowers wafer-sort yield. The move to just-in-time manufacturing (see [Chap. 15](#)) is one attempt to increase yields and decrease the manufacturing costs associated with increased in-line inventories.

Wafer-Sort Yield Formulas

The ability to understand and predict wafer-sort yields with some accuracy is essential to the operation of a profitable and reliable chip supplier. Over the years, a number of models have been developed that relate process, defect densities, and chip size parameters to the wafer-sort yield.⁵ Five yield model formulas are shown in [Fig. 6.12](#). Each relates different parameters to the wafer-sort yield. As the chips get larger, the number of process steps increases, the feature size decreases the sensitivity to smaller defect sizes increases, and more of the background defects become killer defects.

Negative binomial

$$Y_r = \frac{1}{\left(1 + \frac{AD_0}{\alpha}\right)^\alpha}$$

Murphy

$$Y_r = \left(\frac{1 - e^{-AD_0}}{AD_0}\right)^2$$

Bose-Einstein

$$Y_r = \frac{1}{(1 + AD_0)^n}$$

Seeds

$$Y_r = e^{-\sqrt{AD_0}}$$

Poisson

$$P(k) = \frac{\lambda^k k e^{-\lambda}}{k!}$$

$$Y_r = P(0) = e^{-\lambda} = e^{-AD_0}$$

$$Y_r = \int_0^\infty F(D) e^{-AD} d$$

FIGURE 6.12 Wafer-sort yield models.

Exponential Model

The exponential relationship ([Fig. 6.12](#)) or Poisson model is the simplest and one of the first yield models⁵ developed. It is applicable to individual process steps and assumes a random distribution of defects (D_0) across the wafer. For multistep analysis, the factor (n) equating to the number of process steps is used ([Fig 6.12](#)). This model generally is used for products that contain over 300 die and MSI circuits of lower densities. Die sizes that are smaller are predicted by the Seeds model.

The exponential, Poisson, and Seeds models all illustrate the primary relationship between die area, defect density, and wafer-sort yield. In these, e is a constant with a value of 2.718.

B. T. Murphy proposed a model using a more sophisticated distribution of defects. The Bose-Einstein model adds the number of process steps (n), while in the *negative binomial* model, there is a cluster factor. It accounts for defect distributions that tend to be “clustered” on the wafer surface rather than simply exhibiting a random distribution. Adopted by the SIA in the ITRS, the cluster factor is assigned a value of 2.⁶

In most yield models, the factor for processing steps (n) is actually the number of patterning steps. Experience has proved that the patterning steps contribute the greatest number of point defects, and therefore have a direct bearing on sort yield.

[Figure 6.13](#) illustrates the different predictions of the various yield models.⁷ No two complex circuits have comparable designs or processes. Processes vary from company to company, as does the basic background defect density. These factors make the development of an accurate universal yield model difficult. Most chip companies have developed their own models that reflect their particular manufacturing process and product designs. The models are all defect-driven. That is, they assume that all of the fab processes are under control and that the defect levels are those built into the process. They do not include major process problems, such as a contaminated tank of process gas.

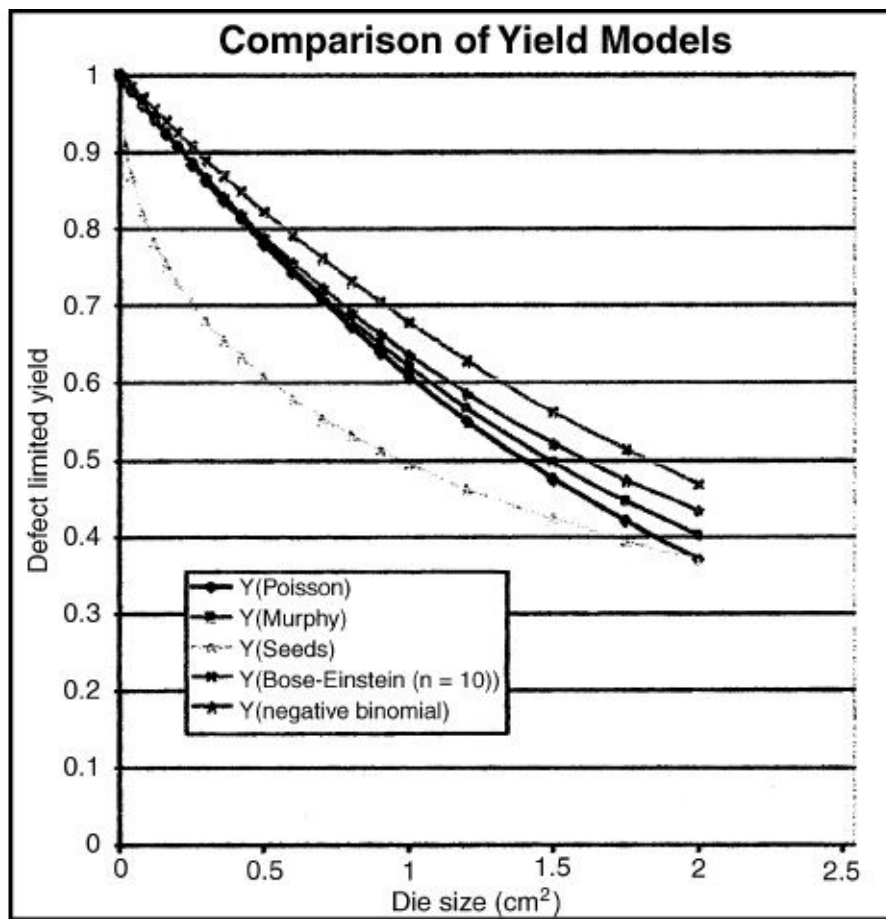


FIGURE 6.13 Yield models showing the die yield as a function of die size defect density.

The defect density used in all the models is not the same as a defect density determined by optical inspection of the wafer surface. The defect density that shows up in the yield models is all-inclusive; it includes contaminants and surface and crystal defects. Further, it predicts only the defects that destroy die: the “killer defects.” Defects that fall in noncritical areas of the chip are not part of the models, nor are situations where two or more defects fall in the same sensitive area.⁸

It is also important to keep in mind that the yield numbers predicted by the formulas are those expected from a process that is basically under control. In reality, the wafer-sort yield will vary from wafer to wafer because of the normal process variations in the fabrication process. A typical wafer-sort yield plot is shown in [Fig. 6.14](#).

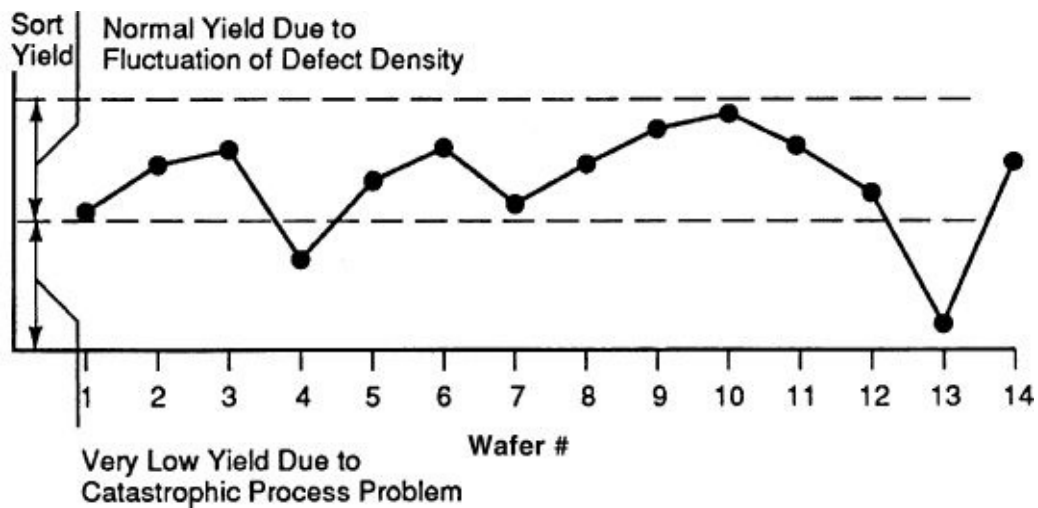


FIGURE 6.14 Plot of wafer-sort yields.

Note that wafer 13 falls far below the normal range of sort yields. In a situation like this, the process engineer would look for some catastrophic process failure such as an out-of-spec layer thickness or a doping layer that is too deep or too shallow.

Assembly and Final Test Yields

After wafer sort, the wafers go to the packaging process, also called *assembly and test*. There, the wafers are separated into the individual die and packaged into a protective enclosure. During this series of steps, there are a number of visual inspections and quality checks of the assembly process.

The last steps of the packaging process is a series of physical, environmental, and electrical tests, known collectively as the *final tests*. (The details of the processes, inspections, and final tests are described in [Chap. 18](#).) After the final tests, the third major yield is calculated, which is the ratio of die passing the final tests compared with the number of *good die* that entered packaging after passing the wafer-sort test.

Overall Process Yields

The overall process yield is the mathematical product of the three major yield points ([Fig. 6.15](#)). This number, expressed in percent, gives the percentage of shipped die as compared with the number of whole die on the starting wafer. It is an inclusive measurement of the success of the entire process.

Overall Yield Formula			
(Wafer-Fab Yield)	(Wafer-Sort Yield)	(Packaging Yield)	= Overall Yield
$\frac{\# \text{ Wafers Out}}{\# \text{ Wafers Started}}$	$\times \frac{\# \text{ Functioning (Good) Die}}{\# \text{ Die On Wafer(s)}}$	$\times \frac{\# \text{ Packages Passing Final Test}}{\# \text{ Die Started Into Packaging}}$	= Overall Yield

FIGURE 6.15 Overall yield formula.

Overall yields vary with several major factors. In [Fig. 6.16](#) is a list of typical process yields and their calculated overall yield. In the first two columns are major process factors that influence the individual and overall yields.

Integration Level	Product Maturity	Fab Yield %	Sort Yield %	Assembly Final Test %	Overall Yield %
ULSI	Mature	95	85	97	78
ULSI	Mid	88	65	92	53
ULSI	Introduction	65	35	70	16
LSI	Mature	98	95	98	91
Discrete	Mature	99	97	98	94

FIGURE 6.16 Typical yields for various products.

First is the integration level of the particular circuit. The more highly integrated the circuit, the lower the expected yield in all categories. Higher integration levels assume a corresponding decrease in feature size. Column 2 lists the maturity of the manufacturing process. Process yields almost always follow an S curve pattern ([Fig. 6.17](#)) through the lifetime of the product in manufacturing. In the beginning, the yield rises rather slowly as the initial bugs are worked out of the process. This is followed by a period when the yields rise rapidly, eventually leveling off as the

limits imposed by the process maturity die size, integration level, circuit density, and defect density. As the table in [Fig. 6.16](#) shows, overall yields can vary from very low (maybe even zero for new or poorly designed products) to the 90 percent range for simpler and mature products. Semiconductor producers consider their yield performance very proprietary, since profit and production control are a direct function of the process yields.

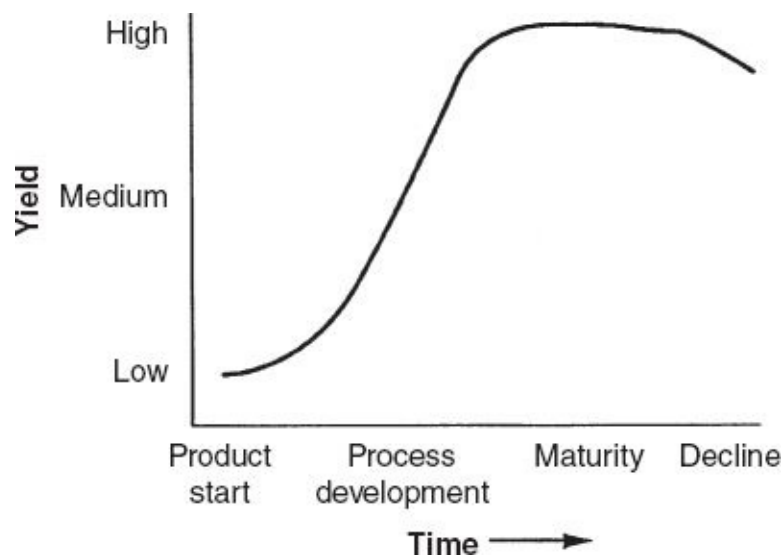


FIGURE 6.17 Yield changes with process maturity levels.

An examination of the yield values in the table reveals that wafer sort is the lowest of the three yield points. This fact illustrates why yield-improvement programs are directed at the many factors that influence the wafer-sort yield. At one time, the improvement of wafer-sort yields had the biggest impact factor on productivity. The advent of larger and more complex chips (such as the megabit memories) has shifted productivity improvements to include other factors, including the cost of ownership of equipment (see [Chap. 15](#)). Successful competition in the megachip era will require wafer-sort yields 90 percent or higher.⁹

Review Topics

Upon the completion of this chapter, you should be able to:

1. Name the three major yield measurement points in the process.
 2. Explain the effect of wafer diameter, die size, die density, number of edge die, and defect density on the wafer-sort yield.
 3. Calculate the accumulative fabrication yield from a list of individual process step yields.
 4. Explain and calculate the overall process yield.
 5. Explain the four major influences on fabrication yield.
 6. Sketch a yield-versus-time curve for different process and circuit maturities.
 7. Explain the relationship between high-process yields and device reliability.
-

References

1. Beaux, L., and Collins, S., "Yield Management," *Handbook of Semiconductor Manufacturing Technology*, 2007, CRC Press, New York, NY: 27-3.
2. Baliga, J., "Yield Management," *Semiconductor International*, Jan. 1998:74.
3. Peters, L., "Speeding the Transition to 0.18 μm ," *Semiconductor International*, Jan. 1998:66.
4. APT Presentation "Overall Roadmap Technology Characteristics," *Industry Strategy Symposium sponsored by The Semiconductor Equipment and Materials Institute*, Jan. 1995.
5. Walker, B., "Motorola VP Defines SubMicron Manufacturing Challenges," *Semiconductor International*, Cahners Publishing, Oct. 1994:21.
6. Ross, R., and Atchison, N., "Yield Modeling," *Handbook of Semiconductor Manufacturing Technology*, 2nd ed., 2008, CRC Press, New York, NY:26-1.
7. Sze, S. M., *VLSI Technology*, 1983, McGraw-Hill Publishing Company, New York, NY:605.
8. Horton, D., "Modeling the Yield of Mixed-Technology Die," *Solid State Technology*, Sep. 1998:109.
9. George, B., and Billatin, S., "Process Control: Covering All of the Bases," *Semiconductor International*, Cahners Publishing, Sep. 1993:80.

CHAPTER 7

Oxidation

Introduction

The ability of a silicon surface to form a silicon dioxide (SiO_2) layer is one of the key factors in silicon technology. This chapter explains the uses, formation, and processes of silicon dioxide growth. Detailed is the all-important tube furnace, which is a mainstay of oxidation, diffusion, heat treatment, and chemical vapor-deposition processes. Other oxidation methods, including rapid thermal processing, are also explained.

Of all the advantages of silicon for the formation of semiconductor devices, the ease of growing a silicon dioxide layer is perhaps the most useful. Whenever a silicon surface is exposed to oxygen, it is converted to silicon dioxide ([Fig. 7.1](#)). Silicon dioxide is composed of one silicon atom and two oxygen atoms (SiO_2). We encounter silicon dioxide daily. It is the chemical composition of ordinary window glass. Semiconductor use requires a purer oxide, which is formed in a highly controlled process. Silicon dioxide layers have been formed with wet and dry thermal oxidation, plasma process, and vapor phase reaction. High-temperature processing in the presence of an oxidant (called *thermal oxidation*) is the process used for semiconductor devices.⁴

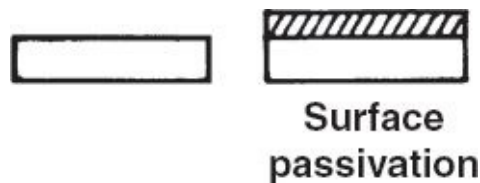


FIGURE 7.1 Surface passivation with silicon dioxide layers.

Although silicon is a semiconducting material, silicon dioxide is a dielectric material. This combination—a dielectric layer formed on a semiconductor—along with other properties of silicon dioxide, makes it one of the most commonly used layers in silicon devices. Silicon dioxide layers are used in devices to pacify the silicon surface, to act as doping barriers and surface dielectrics, and to serve as dielectric parts of device structures. In MOS devices, the most critical layer is the gate, with most gates being formed from a thin layer of silicon dioxide (see [Fig. 7.4](#)).

Silicon Dioxide Layer Uses

Surface Passivation

In [Chap. 4](#), the extreme sensitivity of semiconductor devices to contamination was examined. While a major focus of a semiconductor facility is the control and elimination of contamination, the techniques are not always 100 percent effective. Silicon dioxide layers also play an important role in protecting semiconductor devices from contamination.

Silicon dioxide performs this role in various ways. First is the physical protection of the surface and underlying devices. Silicon dioxide layers formed on the surface are very dense (nonporous) and very hard. Thus, a silicon dioxide layer ([Fig. 7.1](#)) acts as a contamination barrier by physically preventing dirt in the processing environment from getting to the sensitive wafer surface. Also, the hardness of the layer protects the wafer surface from scratches and abuse endured from the outside by the wafer in the fabrication processes.

Another way silicon dioxide protects devices is chemical in nature. Regardless of the cleanliness of the processing environment, some electrically active contaminants (i.e., *mobile ionic contaminants*) end up in or on the wafer surface. During the oxidation process, the top layer of silicon is converted to silicon dioxide. Contaminants on the surface end up in the new layer of oxide, away from the electrically active surface. Other contaminants are drawn up into the silicon dioxide film, where they are less harmful to the devices. In the early days of MOS device processing, it was common to oxidize the wafers and then remove the oxide before further processing, to rid the surface of unwanted mobile ionic contamination.

Doping Barrier

In [Chap. 5](#), doping was identified as one of the four basic fabrication operations. Doping requires creating holes in a surface layer through which specific dopants are introduced into the exposed wafer surface via diffusion or ion implantation. In silicon technology, the surface layer is most often silicon dioxide ([Fig. 7.2](#)). The silicon dioxide left on the wafer acts to block the dopant from reaching the silicon surface. All of the dopants used in silicon technology have a very slow rate of movement in silicon dioxide as compared to their movement in silicon. While the dopants penetrate to the required depth in the exposed silicon, they penetrate only a comparatively short distance into the silicon dioxide surface. It takes only a relatively thin silicon dioxide layer to block the dopants from reaching the silicon surface.

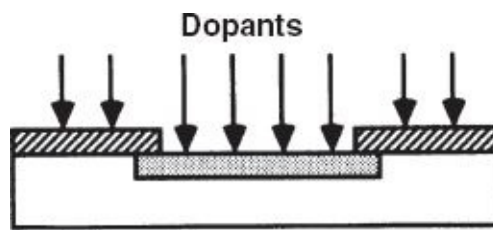


FIGURE 7.2 Silicon dioxide layer as dopant barrier.

Another factor favoring the use of silicon dioxide is a coefficient of thermal expansion similar to that of silicon. In the high-temperature processes of oxidation, diffusion doping, and others, the wafer expands and contracts as it is heated and cooled. The silicon dioxide expands and contracts at close to the same rate as silicon, which means that the wafer will not warp during the heating and cooling.

Surface Dielectric

Silicon dioxide is classified as a dielectric. This means that, under normal circumstances, it does not conduct electricity. When dielectrics are used in electrical circuits or devices, they are referred to as *insulators*. Acting as an insulator is an important role of silicon dioxide layers. [Figure 7.3](#) shows a cross-section of a wafer, with a conductive layer of metal on top of a layer of silicon dioxide. The oxide prevents shorting between the metal layer and the underlying metal just as the insulation on an electric cord prevents the wires from shorting. In this capacity, the oxide must be continuous, that is, have no holes or voids.

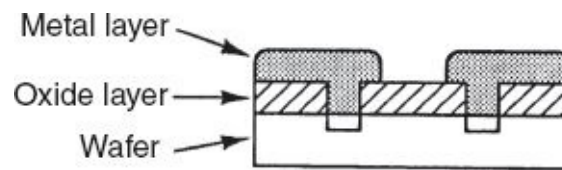


FIGURE 7.3 Oxide layer used as dielectric layer between wafer and metal.

The oxide must also be thick enough to prevent a phenomenon known as *induction*. Induction can occur when the separating layer of oxide is thin enough to allow an electrical charge in a metal layer to cause a buildup of charge in the wafer surface. The surface charge can cause shorting and other unwanted electrical effects. A thick enough layer will prevent an induced charge in the wafer surface. Most of the wafer surface is covered with an oxide layer thick enough to prevent induction from the metal layers. This is called the *field oxide*.

Device Dielectric (MOS Gates)

At the other end of the induction phenomenon is MOS technology. In an MOS transistor, a thin layer of silicon dioxide is grown in the gate region ([Fig. 7.4](#)). Without a charge on the gate, no current passes between the source and drain ([Chap. 16](#)). But with the correct charge, the area under the gate gets an induced charge that allows current flow between the source and drain. The oxide functions as a dielectric whose thickness is chosen *specifically* to allow induction of a charge in the gate region under the oxide (see [Chap. 16](#)). The dominance of MOS technology for ultra-large-scale integrated (ULSI) circuits has made the formation of gate regions a prime focus of process development and concern. Often the thin oxide layer in the gate region will have other dielectric material deposited on top. These combinations, called *gate stacks*, have different dielectric constants that change the electrical function of the transistor. Thermally grown oxides are also used as the

dielectric layer in *capacitors* formed between the silicon wafer and a surface conduction layer ([Fig. 7.5](#)).

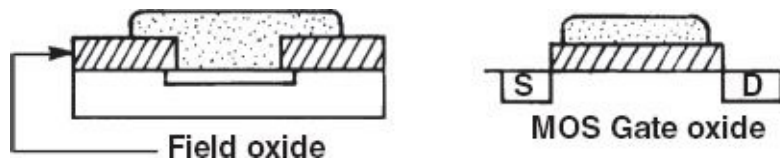


FIGURE 7.4 Silicon dioxide as field oxide and in MOS gate.

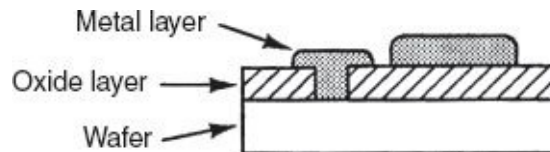


FIGURE 7.5 Silicon dioxide layer in solid-state capacitor.

Silicon dioxide dielectric layers are also used as insulating layers between metal layers in multilayer device structures. In this application, the silicon dioxide layers are deposited with *chemical vapor deposition* (CVD) techniques rather than thermal oxidation (see [Chap. 12](#)).

Device Oxide Thicknesses

The silicon dioxide layers used in silicon-based devices vary in thickness. At the thin end of the scale are advanced MOS gate oxides. Technical advances have allowed gate thickness down to 1 nm (10 angstroms).² At the thick end are field oxides. [Figure 7.6](#) lists the thickness ranges for the major uses.

Silicone Dioxide Thickness, Å	Application
10–100	Tunneling gates
150–500	Gate oxides, capacitor dielectrics
200–500	LOCOS pad oxide
2000–5000	Masking oxides, surface passivation
3000–10,000	Field oxides

FIGURE 7.6 Silicon dioxide thickness chart.

Thermal Oxidation Mechanisms

Thermal oxide growth is a simple chemical reaction, as shown in [Fig. 7.7](#). This reaction takes place even at room temperature. However, an elevated temperature is required to achieve quality oxides in reasonable process times for practical use in circuits and devices. Oxidation temperatures are between 900° and 1200°C.

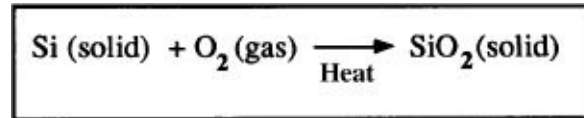


FIGURE 7.7 Reaction of silicon and oxygen to form silicon dioxide.

Although the formula shows the reaction of silicon with oxygen, it does not illustrate the growth mechanism of the oxide. To understand the growth mechanism, consider a wafer placed in a heated chamber and exposed to oxygen gas ([Fig. 7.8a](#)). Initially, the oxygen atoms combine readily with the silicon atoms. This stage is called *linear* because the oxide grows in equal amounts for each unit of time ([Fig. 7.8b](#)). After approximately 1000 angstroms (Å) of oxide is grown, a limit is imposed on the linear growth rate. [An angstrom is one ten-thousandth of a micron (μm); in other words there are 10,000 Å in 1 μm.]

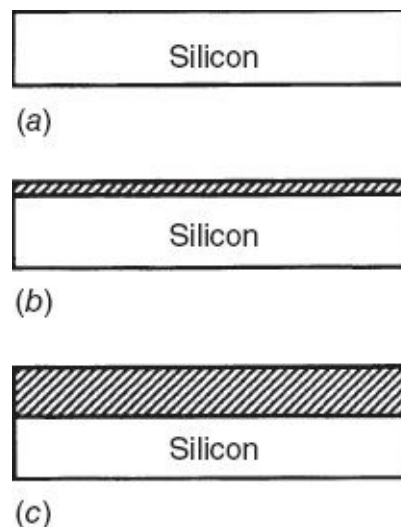


FIGURE 7.8 Silicon dioxide growth states. (a) Initial, (b) linear, and (c) parabolic.

For the oxide layer to keep growing, the oxygen and silicon atoms must come in contact with each other. However, the initially grown layer of silicon dioxide separates the oxygen in the chamber from the silicon atoms of the wafer surface.

For oxide growth to continue, either the silicon in the wafer must migrate through the already grown oxide layer to the oxygen in the vapor, or the oxygen must migrate to the wafer surface. In the thermal growth of silicon dioxide, the oxygen migrates (the technical term is *diffuses*) through the existing oxide layer to the silicon wafer surface. Thus, the layer of silicon dioxide consumes silicon atoms from the wafer surface—the oxide layer *grows into* the silicon surface.

With each succeeding new growth layer, the diffusing oxygen must move further to reach the wafer. The effect is a slowing of the oxide *growth rate* with time. This stage of oxidation is called the *parabolic stage*. When graphed, the mathematical relationship of the oxide thickness, growth rate, and time takes the shape of a parabola. Other terms used for this second stage of growth are *transport-limited reaction*, or *diffusion-limited reaction*, which means that the growth rate is limited by the transportation (diffusion) of the oxygen through the oxide layer already grown. The linear and parabolic stages of growth are illustrated in [Fig. 7.9](#). The formula in [Fig. 7.10](#) expresses the fundamental parabolic relationship for oxide layers above approximately 1200 Å.

$$X = B/A \ t$$

Linear oxidation of silicon

$$X = \sqrt{B \ t}$$

Parabolic oxidation of silicon

X = oxide thickness
 B = parabolic rate constant
 B/A = linear rate constant
 t = oxidation time

FIGURE 7.9 Linear and parabolic growth of silicon dioxide.

$$R = \frac{X^2}{t}$$

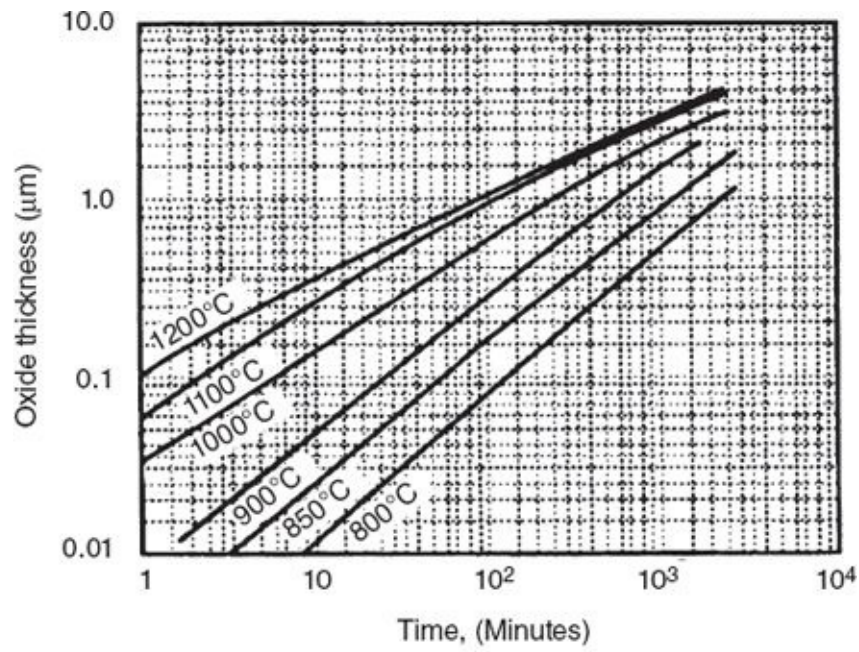
where R = silicon dioxide growth rate
 X = oxide thickness
 t = oxidation time

FIGURE 7.10 Parabolic relationship of SiO₂ growth parameters.

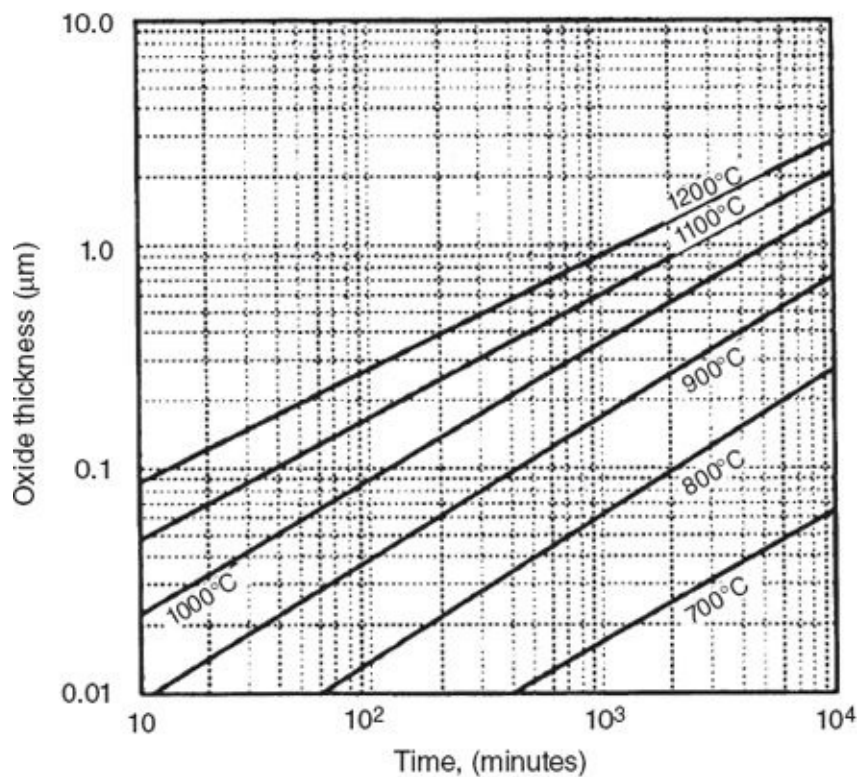
Thus, a growing oxide goes through two stages: the linear stage and the parabolic stage. The change from linear to parabolic is dependent on the oxidizing temperature and other factors (see the following section, “Influences on the

Oxidation Rate”). In general, oxides less than 1000 Å (0.1 μm) are controlled by the linear mechanism. This is the range of most MOS gate oxides.²

The major implication of this parabolic relationship is that thicker oxides require much more time to grow than thinner oxides take. For example, growth of a 2000-Å (0.20-μm) film at 1200°C in dry oxygen requires 6 min ([Fig. 7.11](#)).³ To double the oxide thickness to 4000 Å requires some 220 min—over 36 times as long. This longer oxidation time presents a problem for semiconductor processing. When pure dry oxygen is used as the oxidizing gas, the growth of thick oxide layers requires even longer oxidation times, especially at the lower temperatures. Generally, process engineers want to have the shortest process times possible as are consistent with quality control. The 220 min in the example given is excessive, that is, only one oxidation would be possible in one shift of operation.



(a)



(b)

FIGURE 7.11 Silicon dioxide thickness versus time and temperature in (a) dry oxygen and (b) steam.

One way to achieve faster oxidations is to use water vapor (H_2O) instead of oxygen as the oxidizing gas (oxidant). The growth of silicon dioxide in water vapor proceeds by the reaction shown in [Fig. 7.12](#). In the vapor state, the water is in the form $\text{H}-\text{OH}$. It is composed of one atom of hydrogen (H) and a molecule of

oxygen and hydrogen with a negative charge (OH⁻). This molecule is called the *hydroxyl ion*. The hydroxyl ion diffuses faster than straight oxygen through the oxide layers already on the wafer. The net effect is a faster oxidation of the silicon, as shown in the growth curves in [Fig. 7.11](#).

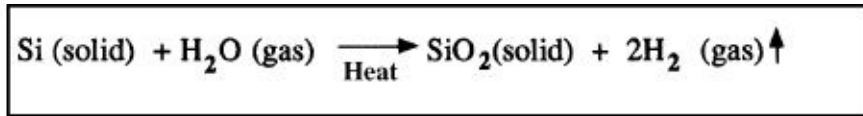


FIGURE 7.12 Reaction of silicon and water vapor to form silicon dioxide and hydrogen gas.

Water vapor at the oxidation temperatures is in the form of steam, and the process is called *steam oxidation*, *wet oxidation*, or *pyrogenic steam*. An oxygen-only oxidation process is called *dry oxidation*. If oxygen only is used, it must be dry (free of any water vapor), or else the type of oxide grown would be that of water vapor and require extra processing.

Notice in the reaction of water vapor and silicon that there are two hydrogen molecules (2H₂) on the right side of the equation. Initially, these hydrogen molecules are trapped in the solid silicon dioxide layer, making the layer less dense than an oxide grown in dry oxygen. However, after a heating of the oxide in an inert atmosphere, such as nitrogen (see the section “Oxidation Processes”), the two oxides become similar in structure and properties.

Influences on the Oxidation Rate

The original oxide thickness versus time curves were determined on $\langle 111 \rangle$ -oriented, undoped wafers.⁴ MOS devices are fabricated in $\langle 100 \rangle$ -oriented wafers and the wafer surfaces are doped. Both of these factors influence the oxidation rate for a particular temperature and oxidant environment. Other factors influencing the oxidation growth are impurities intentionally included in the oxide (such as HCl) and oxidation of polysilicon layers.

Wafer Orientation

The orientation of the wafer has an effect on the oxidation growth rate. For example, $\langle 111 \rangle$ planes have more silicon atoms than $\langle 100 \rangle$ planes. The larger number of atoms allows for a faster oxide growth on $\langle 111 \rangle$ -oriented wafers than for $\langle 100 \rangle$ -oriented wafers. [Figure 7.13](#) shows the growth rates for the two orientations in steam. This difference is seen more in the linear growth stage and at lower temperatures.

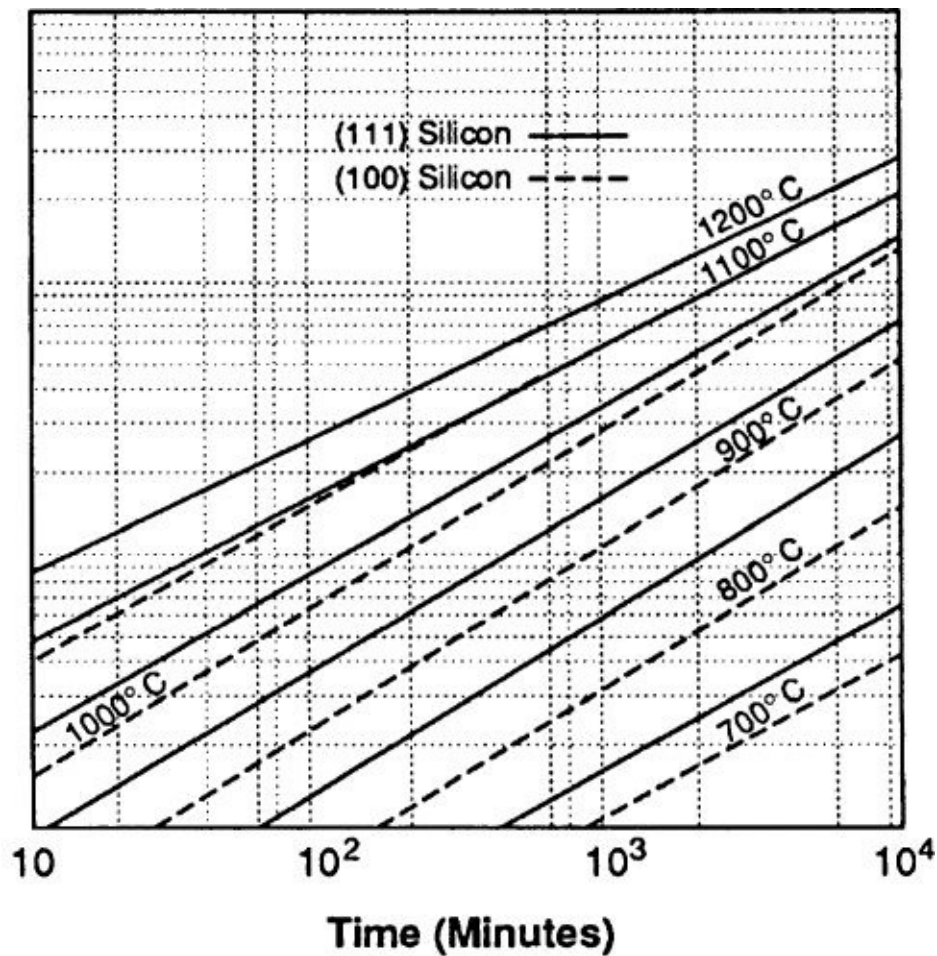


FIGURE 7.13 Oxidation of $\langle 111 \rangle$ and $\langle 100 \rangle$ silicon in steam.

Wafer Dopant Redistribution

The silicon surface being oxidized always has dopants in it. A production silicon wafer starts into the line doped as either an N-type or P-type. Later in the process, wafers have dopant(s) added into the surface from diffusion or ion implant operations. The dopant elements used and their concentration both have effects on the oxidation growth rate. For example, oxides grown over a highly doped phosphorus layer are less dense than those grown over the other silicon dopants. These phosphorus-doped oxides also etch faster and present an etching challenge in the patterning operation due to resist lifting and rapid undercutting.

Another effect on oxidation growth rate is the distribution of the dopant atoms in the silicon after the oxidation is completed.⁵ Recall that during thermal oxidation, the oxide layer grows *down* into the wafer. A question is “What happens to the dopant atoms that were in the layer of silicon converted to silicon dioxide?” The answer depends on the conductivity type of the dopant. The N-type dopants of phosphorus, arsenic, and antimony have a higher solubility in silicon than in silicon dioxide. When the advancing oxide layer reaches them, they move downward into the wafer. The silicon-silicon-dioxide interface acts like a snowplow pushing ahead an ever greater pile of snow. The effect is that there is a higher concentration (called *pile-up*) of N-type dopants at the silicon/dioxide-silicon interface than was originally in the wafer.

When the dopant is the P-type boron, the opposite effect happens. The boron is drawn upward into the silicon dioxide layer, causing the silicon at the interface to be depleted of the original boron atoms (called *depletion*). Both of these effects, pile-up and depletion, have a significant impact on the electrical performance of devices. The exact effects of pile-up and depletion on the dopant concentration profile are illustrated in [Chap. 17](#).

Doping concentration effects on the oxidation rate vary with the dopant type and concentration level. In general, higher-doped regions oxidize faster than more lightly doped regions. Heavily doped phosphorus regions, for example, can oxidize two to five times the undoped oxidation rate.⁶

Doping-induced oxidation effects are more pronounced in the linear stage (thin oxides) of oxidation.

Oxide Impurities

Certain impurities, particularly chlorine from hydrochloric acid (HCl), are included in the oxidizing atmosphere for inclusion in the growing oxide (see “Oxidation Processes”). These impurities have an influence on the growth rate. In the case of HCl, the growth rate can increase from 1 to 5 percent.⁶

Oxidation of Polysilicon

Polysilicon conductors and gates are a feature of most MOS devices or circuits. The device or circuit processes require oxidation of the polysilicon. Compared to the oxidation rates of single-crystal silicon, the rates for polysilicon can be faster, slower, or similar. A number of factors related to the formation of the polysilicon structure influence the subsequent oxidation. They are: the polysilicon deposition method, deposition temperature, deposition pressure, the type and concentration of doping, and the grain structure of the polysilicon.²

Differential Oxidation Rates and Oxide Steps

After the initial oxidation step in a device or circuit fabrication process, the wafer surface has a variety of conditions. Some areas have the field oxide, some are doped, and some are polysilicon regions, and so on. Each of these areas has a different oxidation rate and will increase in oxide thickness depending on the condition. This oxidation thickness difference is called *differential oxidation*. For example, oxidation of an MOS wafer after a polysilicon gate has been formed next to lightly doped source/drain areas ([Fig. 7.14a](#)) results in a thicker oxide growth on gate because silicon dioxide grows faster on polysilicon.

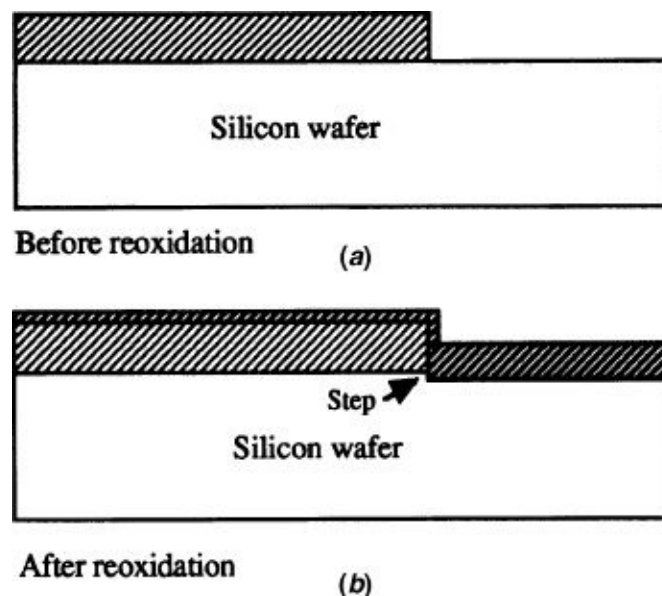


FIGURE 7.14 Differential oxidation of silicon.

Differential oxidation rates cause the formation of steps in the wafer surface ([Fig. 7.14b](#)). Illustrated is a step created by the oxidation of an exposed area next to a relatively thick field oxide. The oxide will grow faster in the exposed area, since

additional oxide growth in the field oxide is limited by the parabolic rate limitation. In the exposed area, the faster-growing oxide will use up more silicon than is used up under the field oxide. The step is shown in [Fig. 7.14b](#).

Thermal Oxidation Methods

The oxide formation reaction formulas include a triangle under the reaction direction arrows. These triangles indicate that the reaction requires energy to proceed. In silicon technology, that energy is usually supplied by heating the wafers and is called *thermal oxidation*. Silicon dioxide layers are grown either at atmospheric pressure or at high pressure. An atmospheric pressure oxidation takes place in a system without intentional pressure control—the pressure is simply that of the atmosphere for the location. There are two atmospheric techniques: *tube furnaces* and *rapid thermal systems* ([Fig. 7.15](#)).

Thermal oxidation		
Atmospheric pressure		Dry oxygen
		Wet oxygen
		Bubbler
		Flash system
		Dry oxidation
	Rapid thermal	Dry oxygen
High pressure	Tube furnace	Dry or wet oxygen
Chemical oxidation		

FIGURE 7.15 Oxidation methods.

Horizontal Tube Furnaces

Horizontal tube furnaces have been the workhorses of tube furnaces since the early 1960s for oxidation, diffusion, heat treating, and various deposition processes. The transition to 200-and 300-mm wafers marked the switch to vertical tube furnaces. Wafer fabs processing smaller-diameter wafers still employ horizontal tube furnaces. The basic tube furnace principles apply to both the systems.

A cross-section of a basic single horizontal three-zone tube furnace is shown in [Fig. 7.16](#). It consists of a long ceramic tube made of mullite (a ceramic), with heating element coils on the inside surface. Separate coils define a zone. Each zone is connected to a separate power supply operated by a proportional band controller. Inside the furnace tube is a quartz reaction tube that serves as the reaction chamber for the oxidation (or other processes). The reaction tube may itself be inside a ceramic liner called a *muffle*. The muffle acts as a heat sink, fostering a more even heat distribution along the quartz tube.

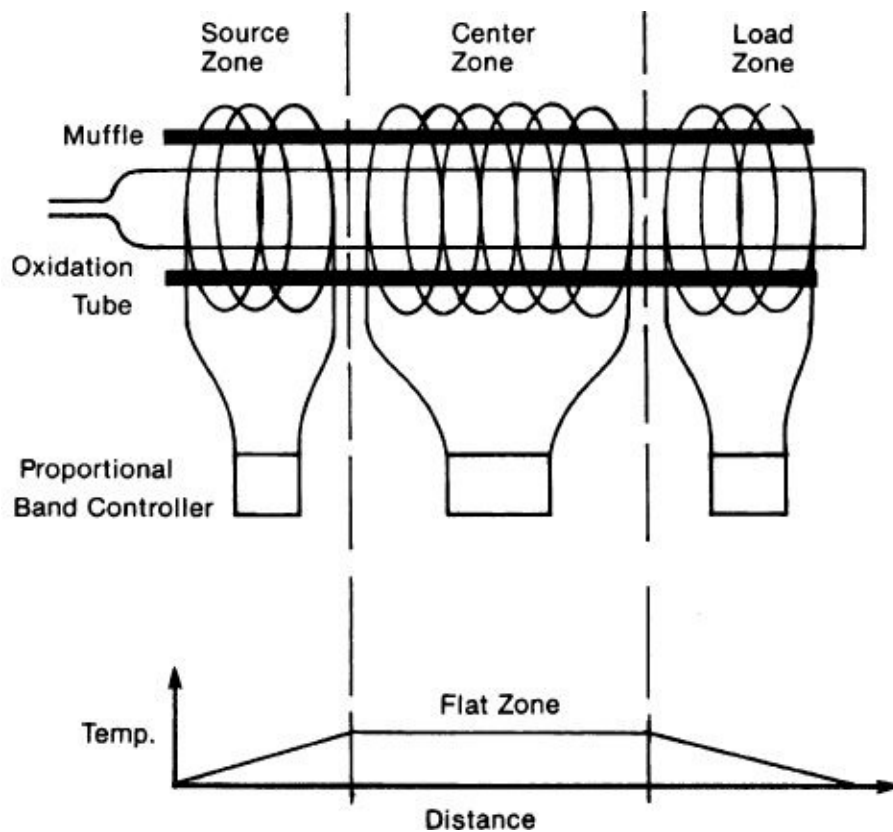


FIGURE 7.16 Cross-section of single horizontal tube furnace with three heating zones.

Thermocouples are positioned against the quartz tube and control power. They send temperature information to the proportional band controllers, which in turn

heat the reaction tube by radiation and conduction. These controllers are very sophisticated and can control temperatures in the center zone (flat zone) to $\pm 0.5^\circ$. For a process that operates at 1000°C , this variation is only ± 0.05 percent. For the oxidation, the wafers are placed on a holder and positioned in the flat zone. The oxidant gas is passed into the tube, where the oxidation takes place.

A production tube furnace is an integrated system of seven various sections:

1. Reaction chamber(s)
2. Temperature control system
3. Furnace section
4. Source cabinet
5. Wafer-cleaning station
6. Wafer load station
7. Process automation

A drawback to horizontal quartz tubes is their tendency to break up and sag at temperatures above 1200°C . The breakup is called *devitrification* and results in small flakes of the quartz tube surface falling onto the wafers.⁸

Temperature Control System

The temperature control system connects thermocouples touching the reaction tube to proportional band controllers that feed the power to the heating coils. Proportional band controllers maintain even temperatures in the tube by feeding in or turning off the current to the coils in proportion to the deviation of the tube temperature from the set point. The closer the tube is to the set-point temperature, the smaller the amount of power that is fed to the coils. This system allows fast recovery of the tube to a cold load without overshoot. Adjustments are made to the controllers until the desired temperatures in the processing section of the tube are achieved.

Overshoot is the raising of the tube temperature too high above the desired process temperature as a result of applying too much power to the coils ([Fig. 7.17](#)).⁹

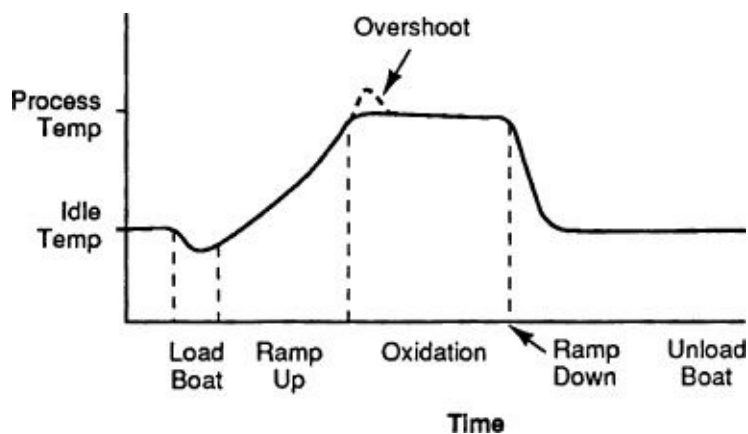


FIGURE 7.17 Temperature levels during oxidation.

The processing of larger-diameter wafers has brought with it the concern of wafer warping ([Fig. 7.18](#)). Silicon expands faster than silicon dioxide. When heated, the silicon dioxide pulls the wafer combination into a concave shape. Wafers that are heated or cooled rapidly will warp to the point of being useless. The degree of warping increases with higher process temperatures—that is, those above 1150°C.



FIGURE 7.18 Wafer warping.

Two methods are employed to minimize warping of wafers in tube furnaces. One is called *ramping* (or *temperature ramping*). Ramping is the procedure of

maintaining the furnace at a temperature several hundred degrees below the process temperature. The wafers are slowly inserted into the furnace at this lower temperature and, after a short stabilization period, the controllers automatically take the furnace up to the process temperature. At the end of the process cycles, the furnace is cooled to the lower temperature before the wafers are removed. During the ramping process, the controllers must maintain the temperature control in the flat zone.

The second antiwarping procedure is the slow loading of the wafer boat into the tube. At loading rates of about 1 in/min, warping is minimized. For large-diameter wafers and large batch sizes, both the methods are used. However, this slow loading extends the total process time. Vertical tube furnaces minimize this problem.

Another requirement of the heating system is a fast recovery time after the wafers are loaded in the tube. A full load of wafers can drop the tube temperature as much as 50°C or more.¹⁰ The heating system works to bring the flat zone to temperature as fast as possible without introducing warping conditions or overshoot. [Figure 7.17](#) illustrates a typical temperature-time recovery curve for a five-zone tube furnace.

A production-level tube furnace will contain three or four tubes (reaction chambers) and a separate temperature control system for each of the tubes. The tubes are arranged vertically above each other in a stack. The tubes open into an exhaust chamber that draws away the spent and heated gases as they exit the tube. This section is connected to the facility's exhaust system, which contains a scrubber to remove toxic gases from the withdrawn gases.

Source Cabinet

Each tube requires a number of gases to accomplish the desired chemical reaction. In the case of oxidation, the gas oxidants of oxygen or water vapor have been detailed. In addition, almost every tube process has nitrogen-flow capability. The nitrogen is used during the loading and unloading stage of the process to prevent accidental oxidation. In the idle condition, nitrogen is kept constantly flowing through the tube. The flow serves to keep dirt out of the system and maintain the preestablished flat zone.

Each of the tube processes requires that the gases be delivered to the tube in a specified sequence, at a specified pressure, at a specific flow rate, and for a specific time. The equipment used to regulate the gases is located in a cabinet attached to the furnace section of the system, and is known as the *source cabinet*. There is a separate unit, called a *gas control panel* or *gas-flow controller*, connected to each tube. The panel consists of solenoids, pressure gauges, mass-flow controller or flowmeters, filters, and timers. In its simplest version, the gas-flow controller consists of manually operated valves and timers. In production systems, the sequencing and timing of the various gases into the tube is controlled by a microprocessor. Required gases are plumbed to the panel. During operation, a timer opens a solenoid to admit the required gas to the panel. Its pressure is controlled by a pressure gauge. Flow amount into the tube is controlled by a flowmeter or mass-flow controller.

A *mass-flowmeter* is preferred in place of a flow meter for its inherent superior control. *Stoichiometric* considerations require that the same amount of material, as measured by its mass, be delivered into the reaction chamber. Semiconductor processes use a thermal-type of mass-flowmeter. The system consists of a heated gas passage tube with two temperature sensors. When no gas is flowing, the temperature sensors are at the same temperature. With the introduction of a gas flow, the downstream sensor reads higher. The difference between the two sensors is related to the amount of heat mass (not the volume) that has moved downstream. The meter has a feedback mechanism to control the gas flow such that a steady amount of material flows through the meter. The source section also contains the microprocessor-controlled valves that meter the gas into the reaction chamber in the right sequence for the right amount of time. A general schematic of a mass-flowmeter is shown in [Fig. 7.19](#). Mass-flowmeters can be set for specific amounts and on-board sensors measure and control the outputs with a feedback control system.¹¹ The piping material used in gas-flow controllers is stainless steel to maintain high levels of cleanliness and to minimize chemical reactions between the gas and tube material.

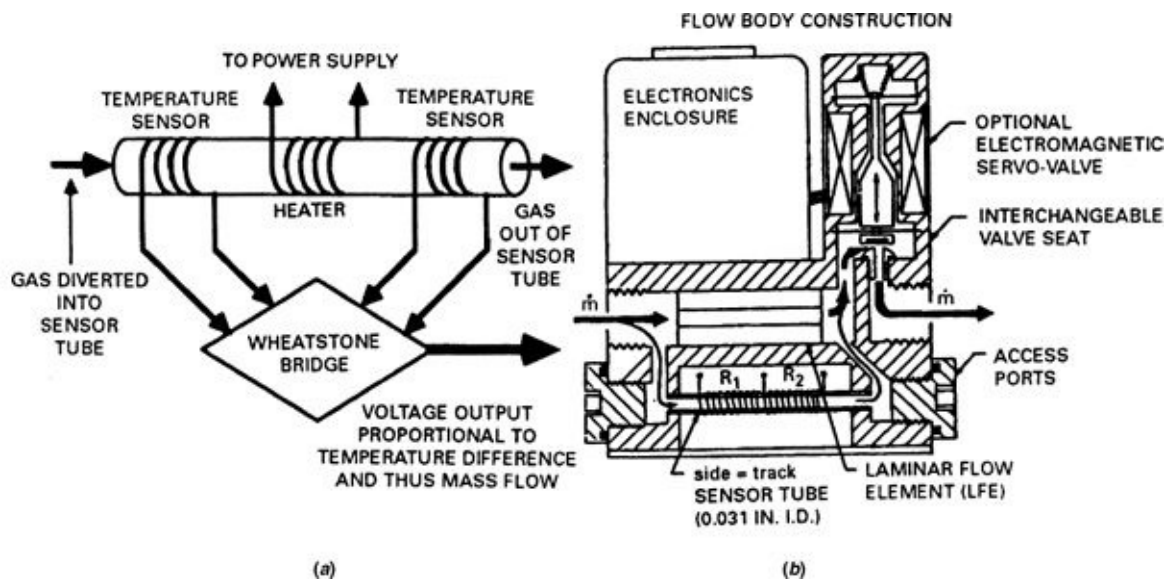


FIGURE 7.19 Mass-flowmeter.

Gases are supplied to the gas-flow controller through piping from the liquid gas supplies in the pad section of the facility, or by smaller *lecture bottles* of gas located at the process tool.

Some processes require a chemical in liquid form. In this situation, a bubbler and liquid source are used. A bubbler consists of a quartz vial designed to admit gas into the liquid. As the gas bubbles through the liquid and mixes with the source vapors in the top of the bubbler, it picks up the source chemicals and carries them into the tube. Bubblers are used in oxidation, diffusion, and CVD processes.

Vertical Tube Furnaces

Horizontal tube furnaces are the oxidation tool of choice for larger-diameter wafers.¹³ There are process problems associated with larger-diameter horizontal tubes. One is keeping the gas streams in a laminar flow pattern in the tube. *Laminar gas flow* is uniform, with no separation of the gases into layers and without turbulence that causes uneven reactions within the tube.

These considerations have resulted in the development of vertical tube furnaces (VTFs), which are the configuration of choice for high-production, large-diameter processes.¹⁴ In this configuration, the tube is held in a vertical position ([Fig. 7.20](#)) with loading taking place from the top or bottom. Tube materials and heating systems are the same as for horizontal systems (see [Fig. 7.21](#)).

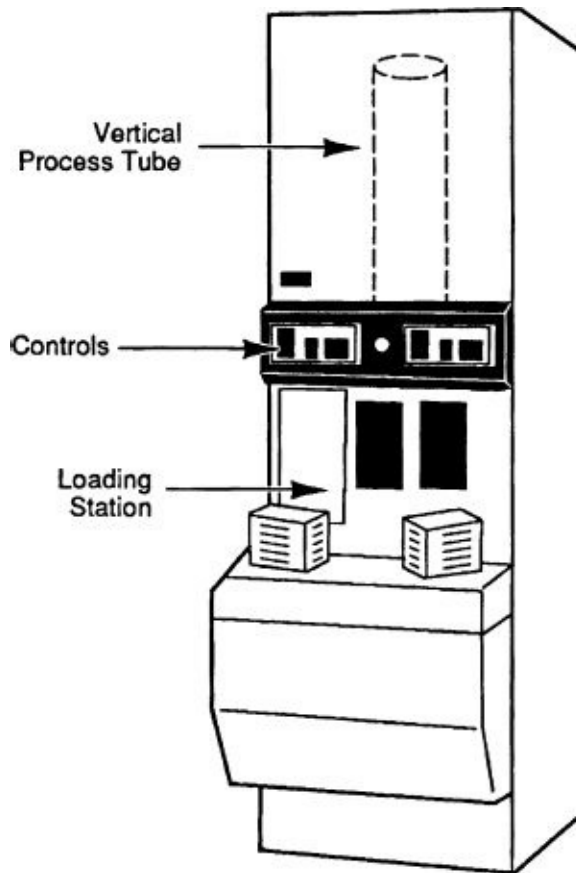


FIGURE 7.20 Vertical tube furnace.

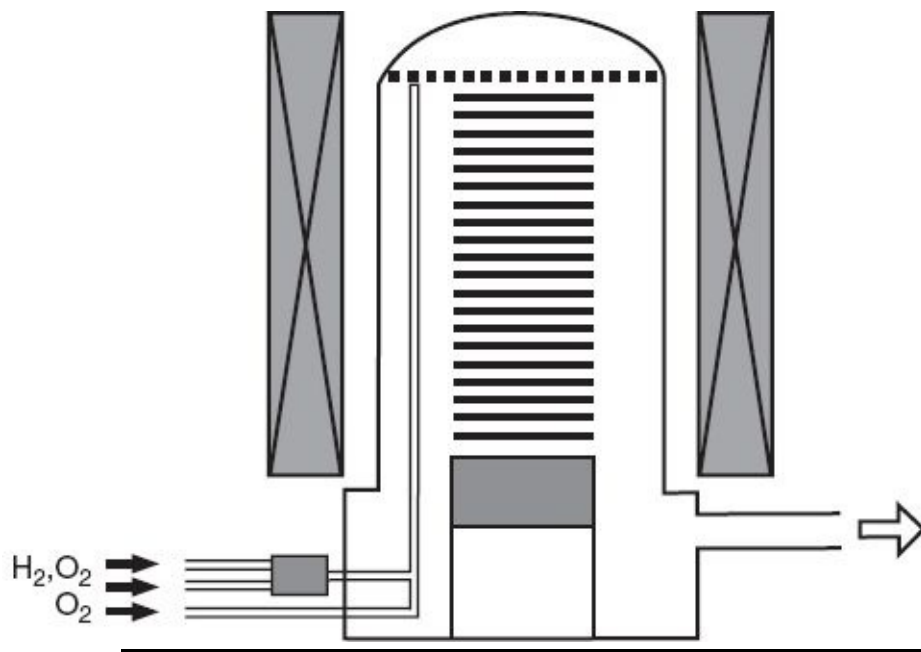


FIGURE 7.21 Cross-section of vertical tube furnace.

The wafers are loaded in standard cassettes and lowered or raised into the flat zone. This action is accomplished without the particulates generated by the cassettes scraping the sides of the tubes. Stacking wafers horizontally allows more wafers per production lot compared to horizontal tube furnaces. An added plus for VTFs is the ease of rotating the wafers in the tube, which produces a more uniform temperature across the wafer. These furnaces have the same subsystems as horizontal furnaces.

Process uniformity is also enhanced by a more uniform (laminar) gas flow in a vertical tube. In a horizontal system, gravity tends to separate mixed gases as they flow down the tube. In a vertical system, the gas moves parallel to gravity, minimizing the gas-separation problem. The boat rotation minimizes gas turbulence. Vertical furnaces are capable of producing 60 percent less process variations than horizontal furnaces produce.¹⁵

Particle generation associated with boat scraping in horizontal systems is virtually eliminated in vertical systems, and the smaller area required for loading results in a cleaner system range.¹⁶

Perhaps one of the most appealing cost aspects of VTFs is the smaller footprint, because these systems are smaller than those of conventional four-stack systems. Vertical systems offer the possibility of locating the furnaces outside the cleanroom with only a load station door opening into the cleanroom. In this arrangement, the cleanroom footprint of the furnace is practically zero, and maintenance can take place from the service chase. Another possible arrangement of vertical furnaces is in an island/cluster configuration. The furnaces are arranged around a central robot that alternately loads several furnaces. A simpler design translates into a more reliable furnace with lower maintenance costs and longer periods of uptime.

Vertical furnaces can be configured to perform any of the oxidation, diffusion, annealing, and deposition processes required in wafer fabrication.

Automation is a huge advantage for VTFs. Loading is by robot of the transfer of the wafers from a FOUP into the furnace cassette in a closed Class-1 environment. Also, ramp-up and ramp-down are faster and the footprint is considerably lower than for horizontal tube stacks.

Advantages of VTFs:

- Large-diameter wafers
- Tighter temperature control (rotation)
- Improved oxide uniformity
- Faster ramp-up and ramp-down
- Cleaner process environment in load-unload station
- Automation compatible

Rapid Thermal Processing

Ion implantation has replaced thermal diffusion due to its inherent doping control. However, ion implantation requires a follow-on heating operation, called *annealing*, to cure out crystal damage induced by the implant process. The annealing step has been traditionally done in a tube furnace. Although the heating anneals out the crystal damage, it also causes the dopant atoms to spread out in the wafer, an undesirable result. This problem led to the investigation of alternate energy sources to achieve the annealing without the spreading of the dopants. The investigations led to the development of rapid thermal process (RTP) technology.

RTP technology is based on the principle of radiation heating ([Fig. 7.22](#)). The wafer is automatically placed in a chamber fitted with gas inlets and exhaust outlets. Inside, a heat source above (and sometimes below) the wafer provides the rapid heating. Heat sources include graphite heaters, microwave, plasma arc, and tungsten halogen lamps.¹⁷ Tungsten halogen lamps are the most popular.¹⁸ The radiation from the heat source couples into the wafer surface and brings it up to the process temperatures of 800°–1050°C at rates of 50°–100°C per second.¹⁹ The same temperature would take minutes to reach in a conventional tube furnace. A typical time-temperature cycle is shown in [Fig. 7.23](#). Likewise, cooling takes place in seconds. With radiation heating, because of its very short heating times, the body of the wafer never comes up to temperature. For the ion-implant annealing step, this means that the crystal damage is annealed while the implanted atoms stay in their original location.

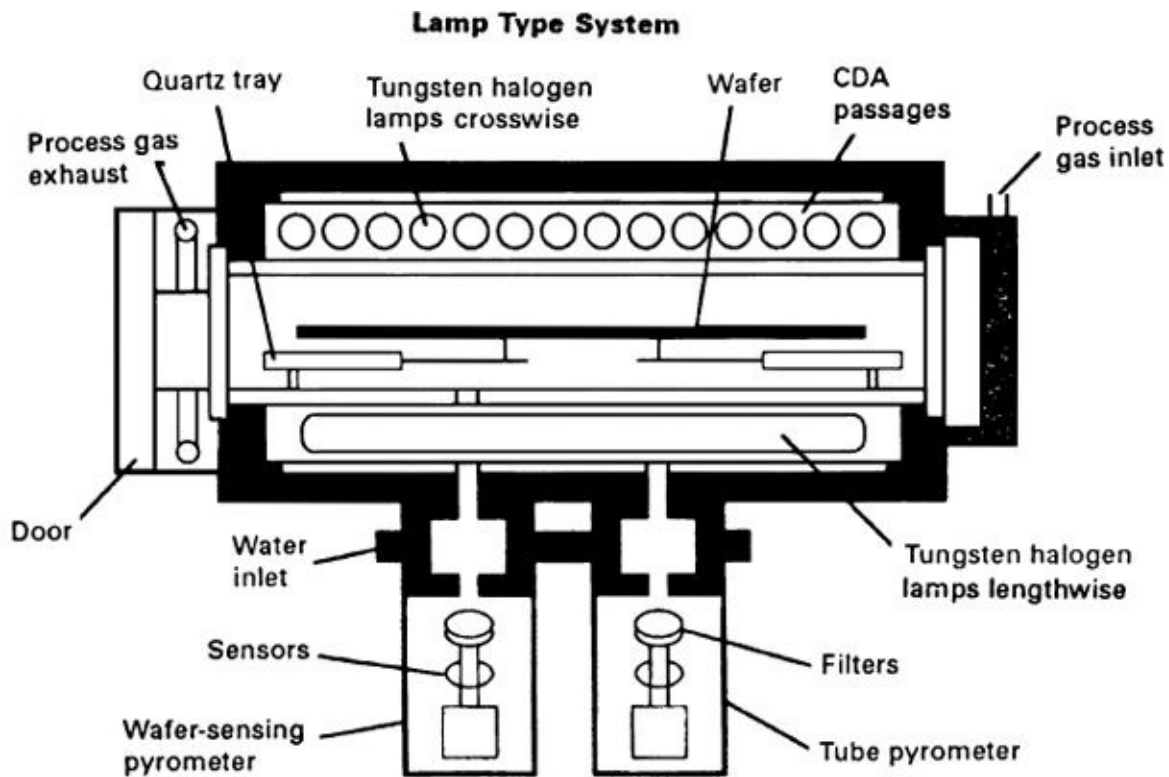


FIGURE 7.22 RTP design. (Source: Semiconductor International, May 1993.)

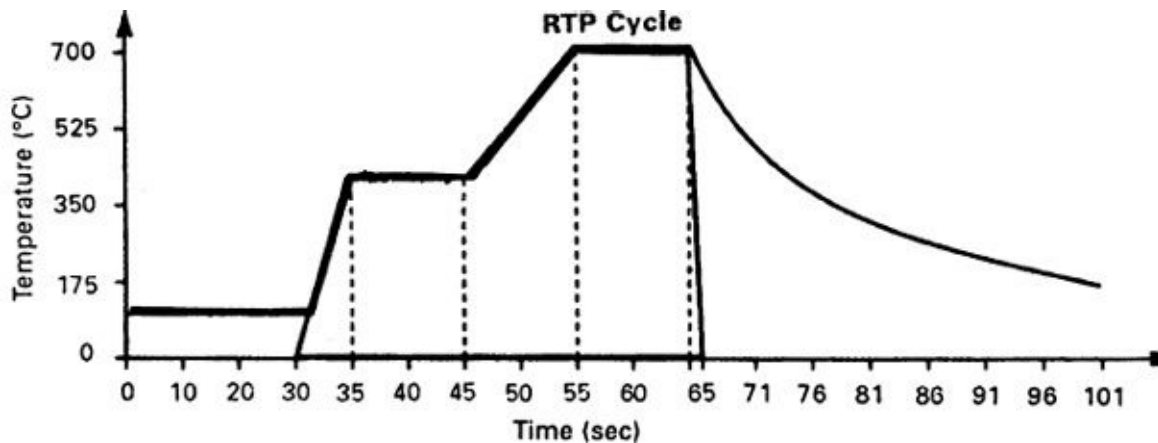


FIGURE 7.23 Example of RTP time/temperature curve. (Source: Semiconductor International.)

Use of RTP reduces the *thermal budget* required for a process. Every time a wafer is heated near diffusion temperatures, the doped regions in the wafer continue to spread down and sideways (see [Chap. 11](#)). Every time a wafer is heated and cooled, more crystal dislocations form (see [Chap. 3](#)). Thus, minimizing the total time a wafer is heated allows more dense designs and fewer failures from dislocations.

Another advantage is single-wafer processing. The move to larger-diameter wafers has introduced uniformity requirements that in many processes are best met in a single-wafer process tool.

RTP technology is a natural choice for the growth of thin oxides used in MOS

gates. The trend to smaller feature sizes on the wafer surface has brought along with it a decrease in the thickness of layers added to the wafer. Layers undergoing dramatic reduction in thickness are thermally grown gate oxides. Advanced production devices are requiring gate oxides in the 10 Å range. Oxides this thin are hard to control in conventional tube furnaces due to the problem of quickly supplying and removing the oxygen from the system. RTP systems can offer the needed control by their ability to heat and cool the wafer temperature very rapidly. RTP systems used for oxidation, called *rapid thermal oxidation* (RTO) systems, are similar to the annealing systems but have an oxygen atmosphere instead of an inert gas. A typical time-temperature-thickness relationship for RTO is shown in [Fig. 7.24](#).

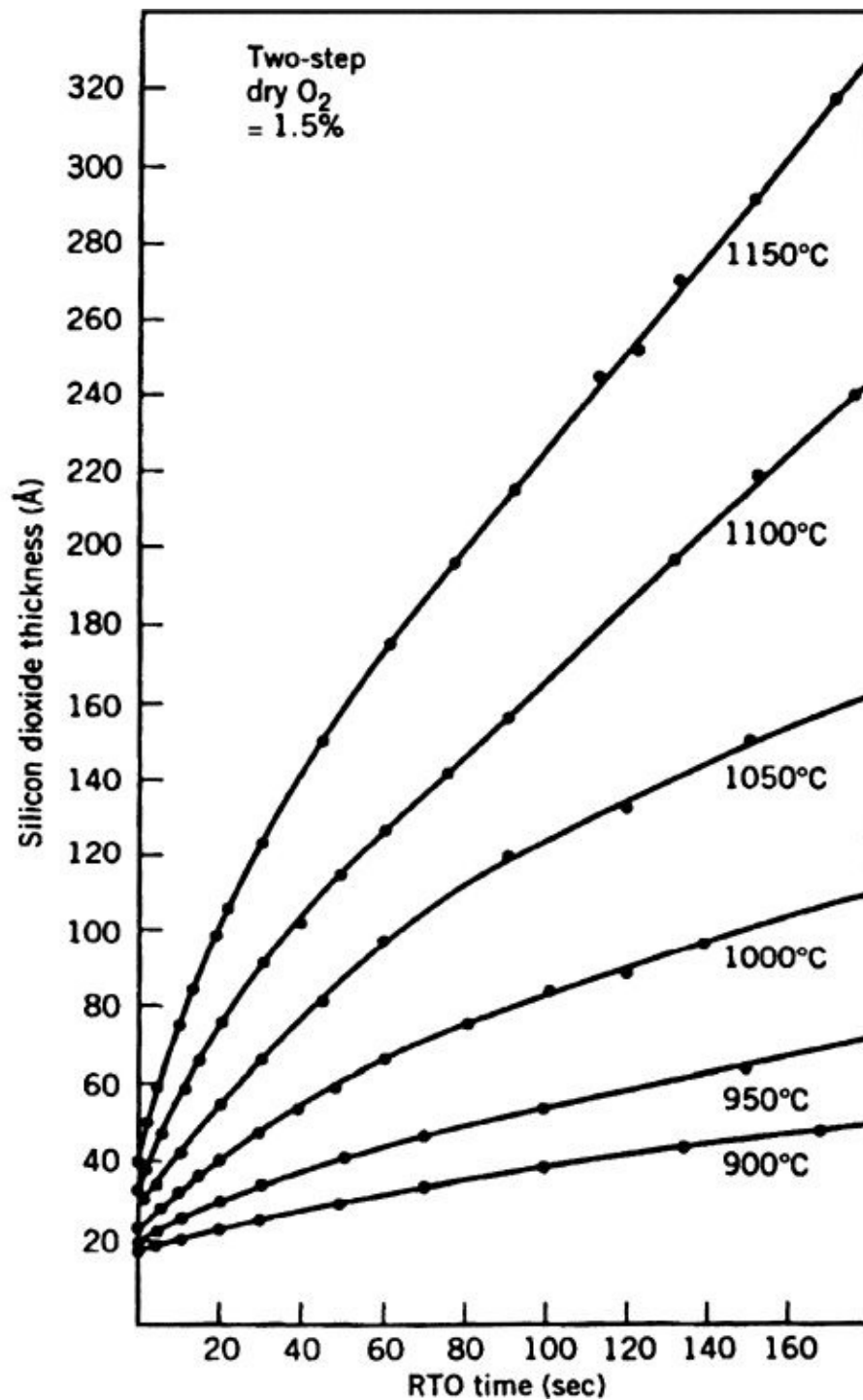


FIGURE 7.24 Oxidation of silicon by RTO. (Source: Ghandhi, VLSI Fabrication Principles.)

Other processes using RTP technology include wet oxide (steam) growth, localized oxide growth, source or drain activation after ion implant, LPCVD polysilicon, amorphous silicon, tungsten, salicide contacts, LPCVD nitride, and LPCVD oxide.²⁰ RTP systems come in atmospheric, low-pressure, and ultra-high-vacuum designs.

Temperature control across a wafer is different in a radiation chamber from that

in a furnace tube. In an RTP system, the wafer never comes to thermal stability. The problem is particularly acute at the wafer edges. Another problem comes from the number and different layers already on an in-process wafer. These different layers each absorb the heating radiation in a different way, resulting in temperature differences across the wafer, which in turn contribute to temperature nonuniformity. This phenomenon is called *emissivity* and is a property of the particular material and the wavelength of the heating radiation. Temperature nonuniformity creates nonuniform process results in and on the wafer surface, and if the temperature differential is high enough, crystal slip at the wafer's edge.

Solutions to the problem include lamp placement and control of individual lamps in the system along with top and bottom lamps. Some systems have a heated annular ring to keep the edge of the wafer within the required temperature range. Process temperatures are usually measured by thermocouples; however, they require back contact with the heated wafer, which is impractical in a single-wafer system, and thermocouples have a response time that is longer than some RTP heating cycles. Optical pyrometers are preferred to thermocouples, as they gauge temperatures by measuring characteristic energies given off by the heated object. However, they too are prone to errors, especially on wafers with a number of layers. The difficulty is relating the emission given off by the wafer to the actual temperature on the surface. Solutions to this problem include a backside seal layer of silicon nitride to minimize backside emissivity variations²¹ and open-loop lamp control. *Open-loop control* is based on converting lamp control to direct current (dc) to get away from voltage variations to the lamps. Other approaches involve elaborate sampling of and/or filtering of the radiation coming off the wafer to more closely relate the measurement to the wafer-surface temperature. It has also been suggested that measurement of the wafer expansion, which is directly due to temperature increases, may be a more reliable and direct measurement technique.²² Given the benefits of RTP, including the ease of automation, it has become a staple of processing.

High-Pressure Oxidation

The thermal budget problem was an impetus (along with others) for high-pressure oxidation. The growth of dislocations in the bulk of the wafer and the growth of hydrogen-induced dislocations along the edges of openings in layers on the wafer surface²³ are two high-temperature oxidation problems. In the first case, the dislocations cause various device performance problems. In the latter case, surface dislocations cause electrical leakage along the surface, or the degradation of silicon layers grown on the wafer for bipolar circuits.

The growth of dislocations is a function of the temperature at which the wafer is processed and the time it spends at that temperature. A solution to this problem is to

perform thermal oxidation processes at a lower temperature. This solution by itself causes the production problem of longer oxidation times. The solution that addresses both problems is high-pressure oxidation ([Fig. 7.25](#)). These systems are configured like conventional horizontal tube furnaces but with one major exception: the tube is sealed, and the oxidant is pumped into the tube at pressures of 10 to 25 atm (10 to 25 times the pressure of the atmosphere). The containment of the high pressure requires encasing the quartz tube in a stainless-steel jacket to prevent it from cracking.

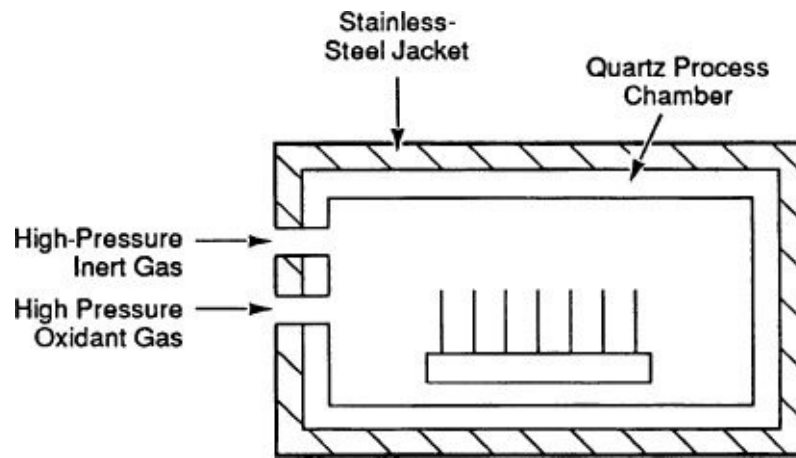


FIGURE 7.25 High-pressure oxidation.

At these pressures, the oxidation proceeds at a faster rate than in atmospheric systems. A rule-of-thumb is that a 1-atm increase in pressure allows a 30°C drop in the temperature. In a high-pressure system, that increase relates to a drop of 300° to 750°C in temperature. This reduction is sufficient to minimize the growth of dislocations in and on the wafers.

Another reason for using high-pressure systems is to maintain the regular process temperature and reduce the time of the oxidation. Other considerations concerning high-pressure systems focus on the safe operation of the system and possible contamination from the additional pumps and piping needed to create the high pressures inside the tube.

Very thin MOS gate oxide growth is a candidate for high-pressure oxidation. The thin oxide must have structural integrity (no holes, and so forth) and have a dielectric strength high enough to prevent charge induction in the gate region. Gate oxides grown in high-pressure processes have higher dielectric strength than similar oxides grown at atmospheric pressure.²⁴ High-pressure oxidation is also a solution for the *bird's beak* problem that occurs during local oxidation of silicon (LOCOS). See the section in [Chap. 16](#), “LOCOS Process.” An unwanted bird’s beak-shaped spur of oxide grows into the active region of an MOS device as in [Fig. 7.26](#). High-pressure oxidation can minimize the bird beak encroachment into the device

area and minimize field oxide thinning during LOCOS processing.²⁵

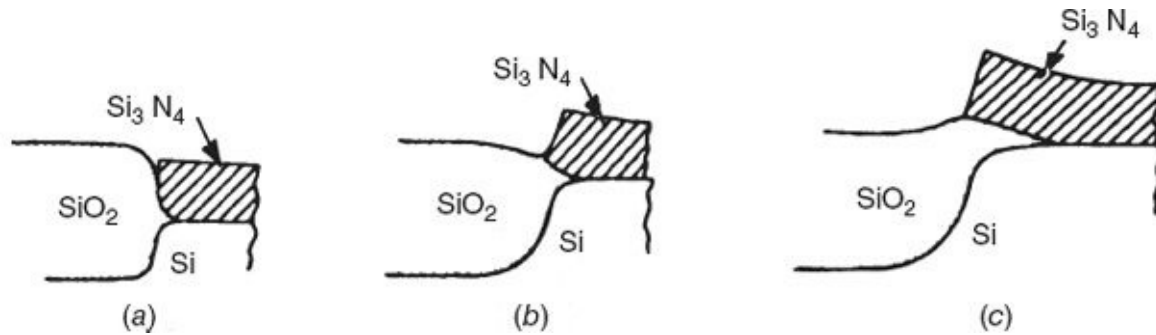


FIGURE 7.26 Bird's beak growth. (a) No pre-etch, (b) 1000 Å pre-etch, and (c) 2000 Å pre-etch. (From Ghandhi, VLSI Fabrication Principles.)

In addition to oxidation, high-pressure systems are finding some use in CVD epitaxial depositions and for flowing glass layers onto the wafer surface.²⁶ Both of these processes are of higher quality when performed at lower temperatures.

Oxidant Sources

Dry Oxygen

When only oxygen is used as the oxidant, it is supplied from the facility source or from tanks of compressed oxygen located in or near the source cabinet. It is imperative that the gas be dry, that is, not contaminated with water vapor. The presence of water vapor in the oxygen would increase the oxidation rate and cause the oxide layer to be out-of-specification (i.e., *out-of-spec*). Dry-oxygen oxidation is the preferred method for growing the very thin ($\approx 1000 \text{ \AA}$) gate oxides required for MOS devices.

Water Vapor Sources

Several methods are used to supply water vapor (steam) into the oxidation tube. The choice of method depends on the level of thickness and cleanliness control required of the oxide layer in the device.

Bubblers

The historic method of creating a steam vapor in the tube has been with a bubbler. It is a quartz vessel with a heater and that holds deionized (DI) water heated close to the boiling point (98° to 99°C), which creates a water vapor in the space above the liquid. A carrier gas carries the vapor into the heated tube where it turns to steam. (An oxidation bubbler is the same construction as a liquid dopant bubbler described in [Chap. 11](#).)

A primary drawback with a bubbler system is that control of the amount of water vapor entering the tube as the water level in the bubbler changes and fluctuates with the water temperature. With bubblers, there is also always concern about contamination of the tube and oxide layer from dirty water or dirty flasks. This contamination potential is heightened by the need to open the system periodically to replenish the water.

Dry Oxidation

New levels of thickness control and cleanliness came with the introduction of MOS devices. The heart of an MOS transistor is the gate structure, and the critical layer in the gate is a thin, thermally grown oxide. Liquid-water-steam systems are unreliable for growing thin, clean gate oxides. The answer was found in the dry oxidation, also called the dryox (or dry steam) process ([Fig. 7.27](#)).

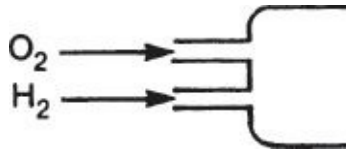


FIGURE 7.27 “Dryox” (dry steam) water vapor source.

In the dry oxidation process and system, gaseous oxygen and hydrogen are introduced directly into the oxidation tube. Inside the tube, the two gases mix and, under the influence of the high temperature, form steam. The result is a wet oxidation in steam. Dryox systems offer improved control and cleanliness over liquid systems. First, gases can be purchased in a very clean and dry state. Second, the amounts going into the tube can be very precisely controlled by the mass-flow controller. Dryox is the preferred general oxidation method for production for all advanced devices.

A drawback to dryox systems is the explosive property of hydrogen. At oxidation temperatures, hydrogen is very explosive. Precautions used to reduce the explosion potential include separate oxygen and hydrogen lines to the tube and flowing excess oxygen into the tube. The excess oxygen ensures that every hydrogen molecule (H_2) will combine with an oxygen atom to form the nonexplosive water molecule, H_2O . Other precautions used are hydrogen alarms and a hot filament in the source cabinet and in the scavenger end of the furnace to immediately burn off any free hydrogen before it can explode.

Chlorine-Added Oxidation

The thinner MOS gate oxides require very clean layers. Improvements in cleanliness and device performance are achieved when chlorine is incorporated into the oxide. The chlorine tends to reduce mobile ionic charges in the oxide layer, reduce structural defects in the oxide and silicon surface, and reduce charges at the oxide-silicon interface. The chlorine comes from the inclusion of anhydrous chlorine (Cl_2), anhydrous hydrogen chloride (HCl), trichloroethylene (TCE), or trichloroethane (TCA) into the dry oxygen gas stream. When the gases chlorine and

hydrogen chloride are used, they are metered into the tube along with the oxygen from separate flowmeters in the gas-flow controller. When the liquid sources TCE and TCA are used, they are carried into the tube as vapors from liquid bubblers. For safety and ease of delivery, TCA is the preferred source of chlorine. The oxidation-chlorine cycle may take place in one step or be preceded or followed by a dry oxidation cycle.

After oxidation, a quick surface inspection of the wafers is normally done with the aid of an ultraviolet (UV) light. These high-intensity light sources allow the operator to see small particles and stains that are not visible to the naked eye. Sometimes, a microscope inspection of the surface is performed.

Automatic Wafer Loading

Once wafers evolved to 200-mm diameters and above, loading and unloading wafers became a challenge for horizontal tube furnace operations. A batch of wafer in a horizontal quartz wafer holder weighs a lot and the loading into a horizontal tube is awkward and inefficient even for robots ([Fig. 7.28](#)). Loading and unloading individual wafers in and out of a vertical stack furnace is much easier. Pick-and-place machines (sometimes called *robots*) pick wafers out of their transfer pod and place them into the quartz furnace boat. A challenge to any wafer boatloading system is the correct placement of test wafers within the load of device wafers as well as “dummy wafers” often placed at the ends (or top and bottom) of a boatload of wafers. These wafers must be *picked* from other boats.

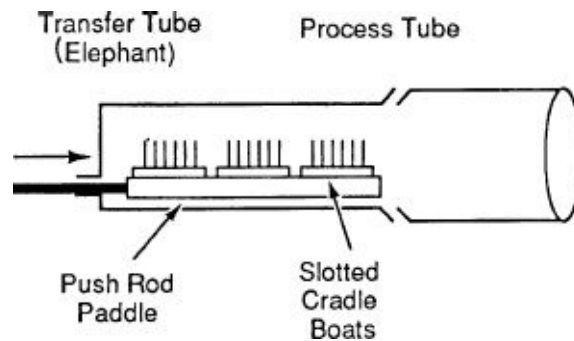


FIGURE 7.28 Transfer tube for loading wafers into a horizontal furnace.

Manual Wafer Handling

Wafers are processed through the cleaning steps in Teflon® or Teflon-derivative wafer holders, also called *boats* or *cassettes*. They are transferred to quartz or silicon carbide holders for the furnace processes.

For these operations with smaller wafers, handling is by vacuum wands or limited-grasp tweezers ([Fig. 7.29](#)). There are also tweezers designed for large-diameter wafers, though the use of tweezers is usually avoided.

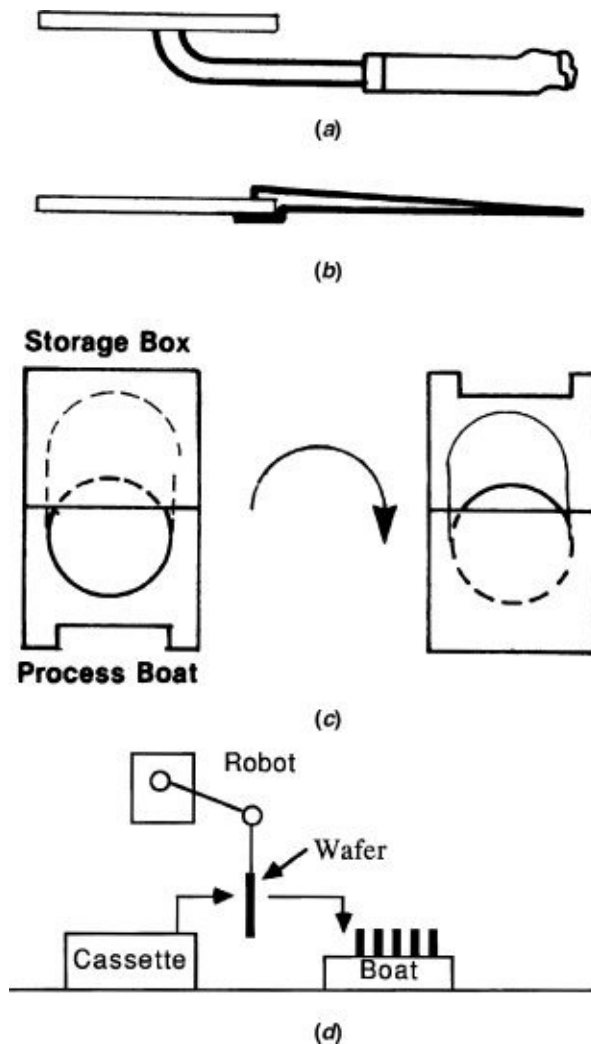


FIGURE 7.29 Wafer handling devices. (a) Vacuum pickup, (b) limited-grasp tweezer, and (c) auto pick-and-place.

Vacuum wands are attached to a vacuum source and are designed to allow grasping of the wafers from the backside. This arrangement minimizes damage and contamination of the sensitive front side of the wafer. Most wafer-handling robots grip the wafers at various points around the perimeter, depending on wafer diameter

and weight.

Oxidation Processes

The general process sequence for oxidation is the same, regardless of the specific oxidation method or equipment used ([Fig 7.30](#)). The wafers are precleaned, cleaned and etched, and loaded in an oxidation boat or oxidation chamber (RTP). The actual oxidation proceeds in different gas cycles ([Fig. 7.31](#)). The first gas cycle occurs as the wafers are being loaded into the tube. Since the wafers are at room temperature and precise oxide thickness is a goal of the operation, the gas metered into the tube during loading is dry nitrogen. The nitrogen is necessary to prevent any oxidation while the wafers are coming up to the required oxidation temperature.

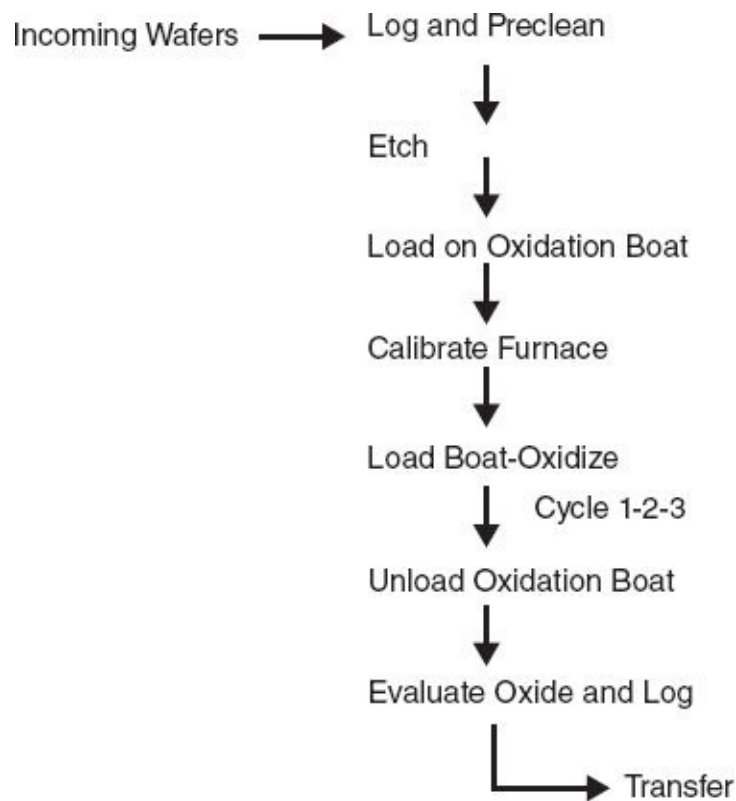


FIGURE 7.30 Oxidation process flow.

Cycle	Gas	Purpose
1.	Nitrogen	Temperature Stabilization in an Inert Atmosphere
2.	Oxygen or Water Vapor	Oxide Growth
3.	Nitrogen	Stop Oxidation and Removal of Wafers in an Inert Atmosphere

FIGURE 7.31 Oxidation process cycles.

Preoxidation Wafer Cleaning

Removal of surface contamination and unwanted native oxides is essential to a successful oxidation process. Contamination can diffuse into the wafer, causing electrical problems in devices and structural integrity problems in the silicon dioxide film. Thin layers of native oxides can alter the thickness and integrity of the grown oxide layer.

Preoxidation processes (see [Chap. 5](#)) typically start with a mechanical scrub, followed by an RCA wet-cleaning sequence to remove organic and inorganic contamination. Finally, an HF or diluted HF etch is performed to rid the surface of native or chemically grown oxides. This is called an HF-last process.

The processing of the wafers through the oxidation process is divided into several distinct steps, as shown in the flow diagram in [Fig. 7.31](#). After the wafers are logged into the station, they receive a thorough cleaning. Wafer cleanliness is essential at all stages of the fabrication process, and especially necessary before any of the operations are performed at high temperature.

Once the wafers are stabilized at the correct temperature, the flow-gas controller switches the gas flow to the selected oxidant. For oxides greater than 1200 Å, the oxidant is usually steam from one of the sources previously discussed. For oxides less than 1200 Å, pure oxygen is usually used because of its greater process control and the cleaner, denser oxide it produces. Thin MOS transistor gate oxides are usually grown in oxygen, at lower temperatures (900°C). Oxidation process times can require hours in the furnace. One alternative is to grow thin gate oxides in wet oxygen environments to reduce process time, but in a reduced pressure. Lowering the pressure maintains oxide density and structural integrity.²⁷ The thinner MOS gate oxides require very clean layers. Improvements in cleanliness and device performance are achieved when chlorine is incorporated into the oxide.²⁸

The oxidation-chlorine cycle may take place in one step or be preceded or followed by a dry oxidation cycle. After the oxidation cycle, the furnace gas is switched back to dry nitrogen. The nitrogen terminates the oxidation of the silicon by diluting and removing the oxidant used. It also prevents any oxidation during the wafer-exit step.

Postoxidation Evaluation

After the wafers are removed from the oxidation boats, they will receive an inspection and multiple evaluations. The nature and number of the evaluations depend on the oxide layer and the precision and cleanliness required of the particular circuit being fabricated. (The details of the evaluations performed are explained in [Chap. 14](#).)

A requirement of the oxidation process is a uniform and contamination-free (or reduced) layer of silicon dioxide on the wafer. As the wafers proceed through the wafer-fabrication operations, there is a buildup of both thermally grown oxides and other deposited layers on the wafer surface. These other layers interfere with the determination of the quality of a particular oxide. For this reason, each batch of oxidized wafers going into a tube includes a number of test wafers (also called monitor wafers), with bare surfaces placed at strategic locations on the wafer boat. Test or monitor wafers are necessary for the evaluations that are destructive or require large undisturbed areas of oxide. On conclusion of the oxidation operation, they are used for the evaluation of the process. Some of the evaluations are performed by the oxidation operator and some are performed off-line in quality control (QC) labs.

Surface Inspection

A quick check of the cleanliness of the oxide is performed by the operator as the wafers are unloaded from the oxidation boat. Each wafer is viewed under a high-intensity UV light. Surface particulates, irregularities, and stains are readily apparent in UV light.

Oxide Thickness

The thickness of the oxide is of major importance. It is measured on test wafers by a number of techniques (see [Chap. 14](#)). The techniques are: (old) color comparison, fringe counting, interference, (new) ellipsometers, stylus apparatus, and scanning electron microscopes (SEMs).

Oxide and Furnace Cleanliness

In addition to the physical contaminants of particles and stains, the oxide should have no more than a minimum number of mobile ionic contaminants. These are detected by the sophisticated capacitance-voltage (C/V) technique, which detects the total number of mobile ionic contaminants present in the oxide. C/V cannot identify the origin of the contaminants, which may come from the process equipment, the gases, the wafers, or the cleaning process. Therefore, C/V evaluation is used as a go/no-go assessment of the wafers and serves as a check of the total furnace operation.

In most fabrication lines, C/V analysis is also used to certify the cleanliness of the furnace and its associated parts. An oxide with a low-mobile ionic contamination level certifies that the entire system is clean. When the oxide fails the test, more investigation is necessary to identify the source.

A second oxide-cleanliness-related parameter is dielectric strength. This parameter measures the dielectric (nonconducting) property of the oxide by using the destructive oxide rupture test.

A third cleanliness factor is the index of refraction of the oxide. Refraction is the property of a transparent substance that causes light to bend as it travels through it. The apparent versus actual location of an object on the bottom of a body of water is an example of refraction. The index of refraction of a pure oxide is 1.46. Variations of this value come about from impurities in the oxide. A constant index of refraction is the starting point for several of the interference thickness-measurement techniques. Variations in the index can lead to erroneous thickness measurements. The index of refraction is measured by interference and ellipsometry techniques (see [Chap. 14](#)).

The various components of an oxidation process may be organized in a *cluster arrangement* (see [Chap. 15](#)).

Thermal Nitridation

An important factor in the production of small high-performance MOS transistors is a thin gate oxide. However, in the 100-Å (and less) range, silicon dioxide films tend to be of poor quality and difficult to control (see [Fig. 7.32](#)). An alternative to silicon dioxide films is a thermally grown silicon nitride (Si_3N_4) film. Si_3N_4 is denser than silicon oxide and has fewer pinholes in these thin ranges. It also is a good diffusion barrier. Growth of thin films using silicon nitride starts with an initial rapid growth rate, but then flattens out, providing greater thickness control of the film. This characteristic is shown ([Fig. 7.32](#)) in the growth of silicon nitride formed by the exposure of the silicon surface to ammonia (NH_3) between 950° and 1200°C.²⁹

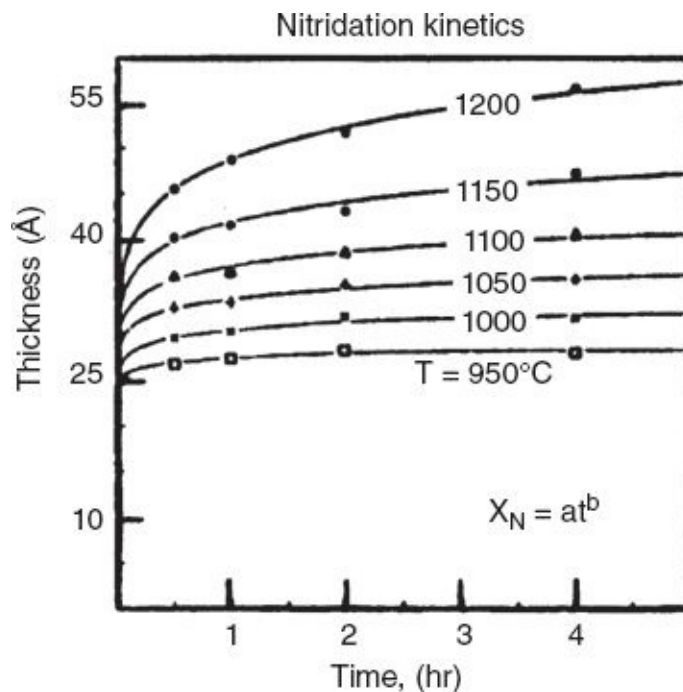


FIGURE 7.32 Nitridation of $\langle 100 \rangle$ silicon. (Source: Wolf, Silicon Process.)

Some advanced devices use silicon oxynitride (SiO_xN_y) films. They are also called *nitrided-oxide* or *nitrooxide* films. These are formed from the nitridation of silicon oxide films. Unlike silicon dioxide films, oxynitride films vary in composition depending on the growth process.³⁰ Another MOS gate structure is a sandwich of oxide/nitride/oxide (ONO).^{31,32}

Review Topics

Upon completion of this chapter, you should be able to:

1. List the three principal uses of a silicon dioxide layer in silicon devices.
 2. Describe the mechanism of thermal oxidation.
 3. Sketch and identify the principal sections of a tube furnace.
 4. List the two oxidants used in thermal oxidation.
 5. Sketch a diagram of a dryox oxidation system.
 6. Draw a flow diagram of a typical oxidation process.
 7. Explain the relationship of process time, pressure, and temperature on the thickness of a thermally grown silicon dioxide layer.
 8. Describe the principles and uses of rapid thermal, high-pressure, and anodic oxidation.
-

References

1. Cleasvelin, C. R., Columbo, L., Nimi, H., and Pas, S., *Oxidation and Gate Dielectrics, Handbook of Semiconductor Manufacturing Technology*, 2nd ed., 2008, CRC Press, New York, NY:9-1.
2. Hu, C., "MOSFET Scaling in the Next Decade and Beyond," *Semiconductor International*, Cahners Publishing, Jun. 1994:105.
3. Wolf, S., and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Sunset Beach, CA:1986.
4. Gise, P., and Blanchard, R., *Modern Semiconductor Fabrication Technology*, 1986, Reston Books, Reston, VA:43.
5. Sze, S. M., *VLSI Technology*, 1983, McGraw-Hill Book Company, New York, NY:137.
6. Gise, P., and Blanchard, R., *Modern Semiconductor Fabrication Technology*, 1986, Reston Books, Reston, VA:46.
7. Ghandhi, S. K., *VLSI Fabrication Principles*, 1994, John Wiley & Sons, Inc., New York, NY:464.
8. Sze, S. M., *VLSI Technology*, 1983, McGraw-Hill Book Company, New York, NY:147.
9. Ibid., p. 159.
10. Hill, M., Helman, D., and Rother, M., "Quartzglass Components and Heavy Metal Contamination," *Solid State Technology*, PennWell Publishing Corp., Mar 1994:49.
11. Maliakal, J., Fisher, D., Jr., and Waugh, A., "Trends in Automated Diffusion Furnace Systems for Large Wafers," *Solid State Technology*, Dec. 1984:107.
12. Murray, C., "Mass Flow Controllers: Assuring Precise Process Gas Flows," *Semiconductor International*, Cahners Publishing, Oak Brook, IL, Oct. 1985:72.
13. Burggraaf, P., "Hands-Off Furnace Systems," *Semiconductor International*, Sep. 1987:78.
14. Burggraaf, P., "Verticals: Leading Edge Furnace Technology," *Semiconductor International*, Cahners Publishing, Sep. 1993:46.
15. Singer, P., "Trends in Vertical Diffusion Furnaces," *Semiconductor International*, Apr. 1986:56.
16. Burggraaf, P., "Verticals: Leading Edge Furnace Technology," *Semiconductor International*, Cahners Publishing, Sep. 1993:46.

- [17.](#) Singer, P., "Furnaces Evolving to Meet Diverse Thermal Processing Needs," *Semiconductor International*, Mar. 1997:86.
- [18.](#) Singer, P., "Rapid Thermal Processing: A Progress Report," *Semiconductor International*, Cahners Publishing, May 1993:68.
- [19.](#) Cleasvelin, C. R., Columbo, L., Nimi, H., et al., *Oxidation and Gate Dielectrics, Handbook of Semiconductor Manufacturing Technology*, 2nd ed., 2008, CRC Press, New York, NY:9-4.
- [20.](#) Leavitt, S., "RTP: On the Edge of Acceptance," *Semiconductor International*, Mar. 1987.
- [21.](#) Moslehi, M., Paranjpe, A., Velo, L., et al., "RTP: Key to Future Semiconductor Fabrication," *Solid State Technology*, PennWell Publishing, May 1994:37.
- [22.](#) Ibid., p. 38.
- [23.](#) Singer, P., "Rapid Thermal Processing: A Progress Report," *Semiconductor International*, Cahners Publishing, May 1993:69.
- [24.](#) Toole, D., and Crabtree, P., "Trends in High-Pressure Oxidation," *Microelectronic Manufacturing and Test*, Oct. 1988:1.
- [25.](#) Ghandhi, S. K., *VLSO Fabrication Principles*, 1994, John Wiley & Sons, Hoboken, NJ:466.
- [26.](#) Kim, S., Emami, A., and Deleonibus, S., "High-Pressure and High-Temperature Furnace Oxidation for Advanced Poly-Buffered LOCOS," *Semiconductor International*, May 1994:64.
- [27.](#) Toole, D., and Crabtree, P., "Trends in High-Pressure Oxidation," *Microelectronic Manufacturing and Test*, Oct. 1988:8.
- [28.](#) Dance, B., "Growth of Ultrathin Silicon Oxides by Wet Oxidation," *Semiconductor International*, Feb. 2002:44.
- [29.](#) Wolf, S., and Tauber, R. N., *Silicon Processing for the VLSI Era*:226.
- [30.](#) Ibid., p. 210.
31. Ghandhi, S. K., *VLSI Fabrication Principles*, 1994, John Wiley & Sons, Inc., New York, NY:484.
32. Singer, P., "Directions in Dielectrics in CMOS and DRAMs," *Semiconductor International*, Cahners Publishing, April 1994:57.

The Ten-Step Patterning Process—Surface Preparation to Exposure

Introduction

Patterning is the series of processes that establishes the shapes, dimensions, and placement of the required physical “parts” (components) of the devices and (ICs) on the wafer surface layers. This chapter presents the first four steps of a basic ten-step photo process and a discussion of photoresist chemistry.

Patterning also referred to as *photolithography*, *photomasking*, *masking*, *oxide removal* (OR), *metal removal* (MR), and *microlithography*. Patterning is one of the most critical operations in semiconductor processing. It is the process that sets the dimensions on the parts of the devices and circuits on the wafer surface. The goal of the operation is twofold:

1. To create, in and on the wafer surface, a pattern with the dimensions established in the design phase of the IC or device ([Fig. 8.1](#)).

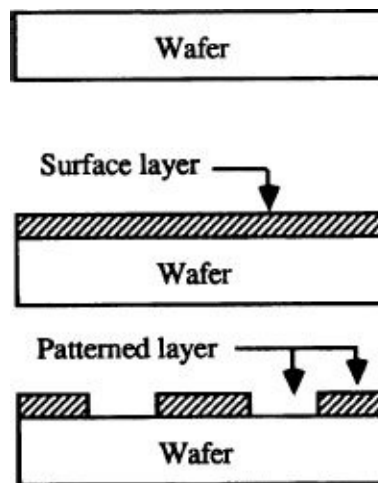


FIGURE 8.1 Basic patterning process.

2. The correct placement of the circuit pattern on the wafer relative to the crystal orientation of the wafer and in a manner that all the layered parts are aligned ([Fig. 8.2](#)).

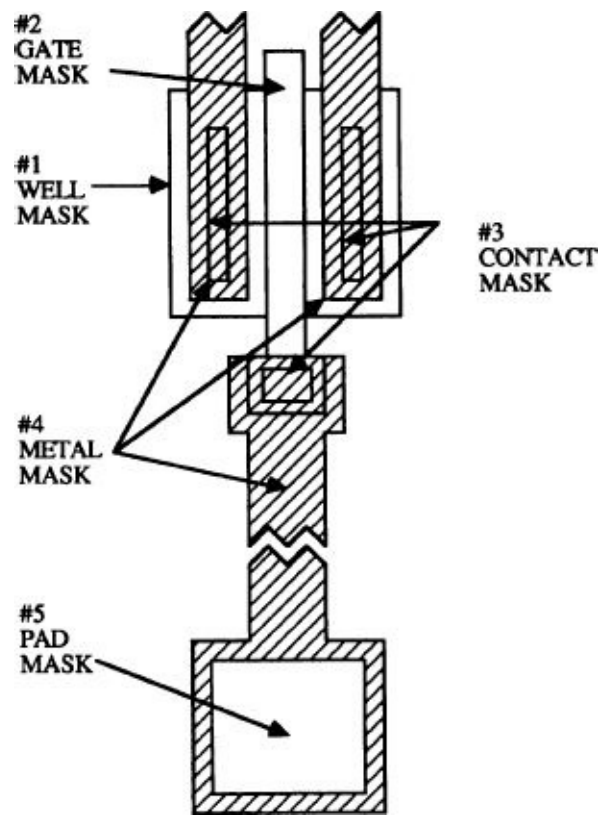


FIGURE 8.2 Five mask set silicon gate transistor.

There are many process variations, but two outcomes. Either a defined part of the wafer surface layer is removed (hole) or defined part of the wafer surface layer is left (island) as shown in [Fig. 8.1](#).

Correct placement is called *alignment* or *registration* of the various circuit patterns and layers. An IC wafer-fabrication process can require 40 or more individual patterning (or masking) steps. This registration requirement is similar to the correct alignment of the different floors of a building. It is easy to visualize that misalignment of elevator shafts and stairwells would render the building useless. In a circuit, the effects of misaligned mask layers can cause the entire circuit to fail.

Additionally, the photolithographic process must control the required dimensions and defect levels. Given the number of steps in each patterning operation and the number of mask layers, the masking process is the chief source of defects. Each masking step in the patterning process contributes variations. A patterning process is one of tradeoffs and balancing explained in detail in [Chap. 10](#).

Overview of the Photomasking Process

Photolithography is a multistep pattern transfer process similar to photography and stenciling. It starts with a circuit design that is translated to the three dimensions of the individual parts of the devices and circuits. Next the X-Y (surface) dimensions, shapes, and alignment of the surface are drawn (composite drawing). Then the composite drawing is separated into the individual masking layers (mask set). This electronic information is loaded into a *pattern generator*. The information from the pattern generator is in turn used to create *reticles* and *photomasks*. Or the information can drive an exposure and alignment tool to directly transfer the pattern to the wafer.

Three major techniques are used to create the individual layer patterns in the wafer layer surface. They are:

1. Replicate the specific layer pattern of a chip on a chrome layer on a quartz plate (reticle). And in turn use the reticle to create a photomask that carries the pattern for an entire wafer (see [Fig. 4.14](#)).
2. Or the reticle can be used to directly pattern the wafer surface layer using a tool called a *stepper* (see [Chap. 10](#)).
3. Or the circuit layer information (dimensions, shapes, alignment, etc.) in the pattern generator can be used to directly guide an e-beam or other source onto the wafer surface (*direct write*) (see [Chap. 10](#)).

The ten-step basic patterning process described here is one that uses a reticle or photomask in the alignment and exposure step. This transfer takes place in two steps. First, the pattern on the reticle or mask is transferred into a layer of photoresist ([Fig. 8.3](#)). Photoresist is a light-sensitive material similar to the coating on old-fashioned photographic film. Exposure to light sources causes changes in its structure and properties. In the example in [Fig. 8.3](#), the photoresist in the region exposed to the light was changed from a soluble condition to an insoluble one. Resists of this type are called *negative acting*, and the chemical change is called *polymerization*. Removing the soluble portion with chemical solvents (developers) leaves a hole in the resist layer that corresponds to the opaque pattern on the mask or reticle.

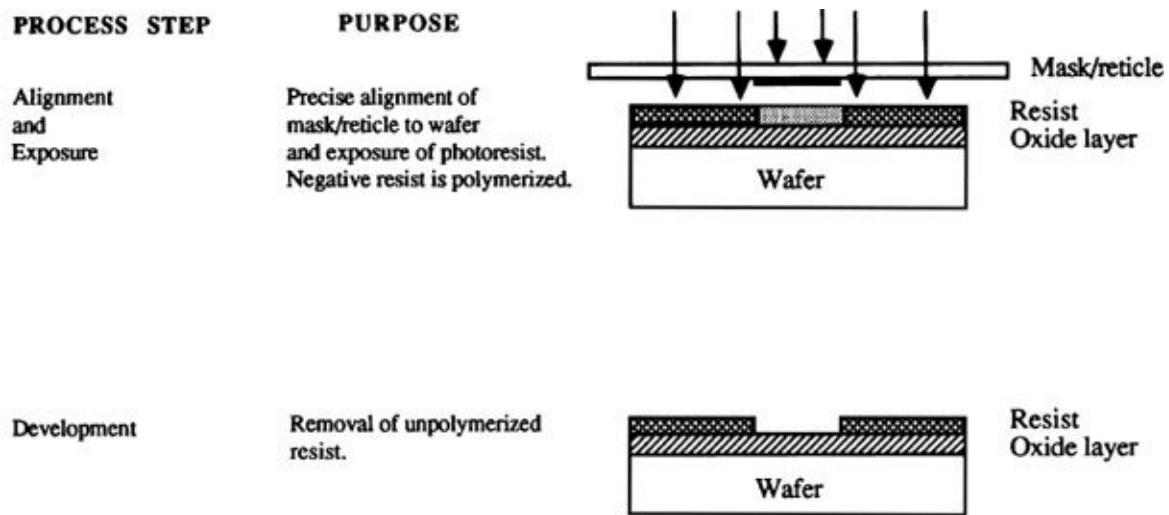


FIGURE 8.3 First-pattern transfer—mask or reticle to resist layer.

The second transfer takes place from the photoresist layer into the wafer surface layer (Fig. 8.4). The transfer occurs when etchants remove the portion of the wafer's top layer that is not covered by the photoresist. The chemistry of photoresists is such that they do not dissolve (or dissolve slowly) in the chemical-etching solutions, that is, they are *etch-resistant*, hence the name *resists* or *photoresists*.

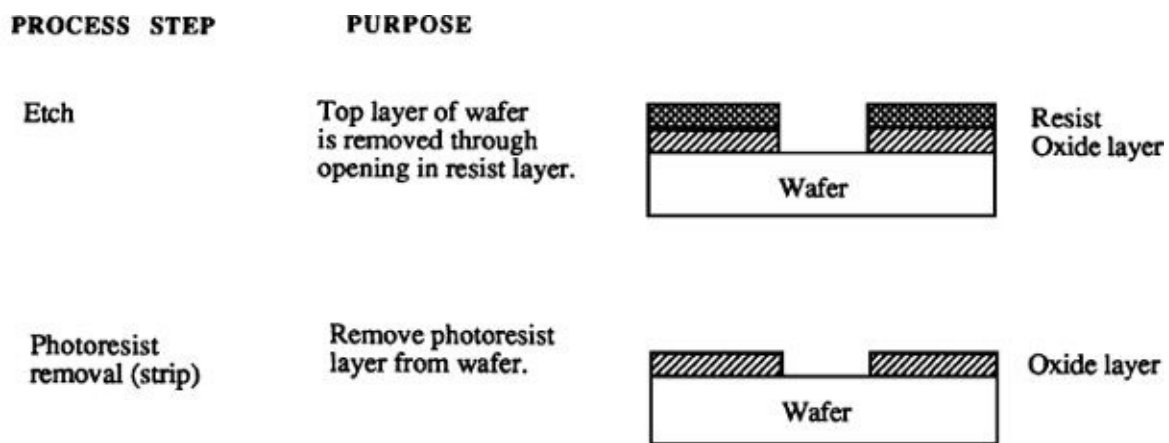


FIGURE 8.4 Second-pattern transfer—resist layer to surface layer.

In the examples shown in Figs. 8.3 and 8.4, the result is a hole etched in the wafer layer. The hole came about because the pattern in the mask was opaque to the exposing light. A mask whose pattern exists in the opaque regions is called a *clear-field mask* (Fig. 8.5). The pattern could also be coded in the mask in the reverse, in a *dark-field mask*. If the same steps were followed, the result of the process would be an island of material left on the wafer surface (Fig. 8.6).

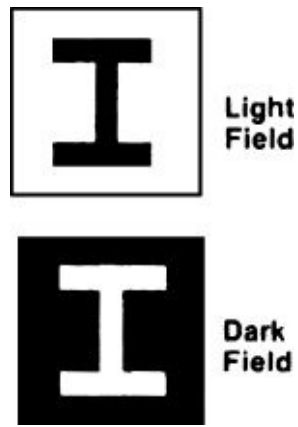


FIGURE 8.5 Mask-reticle polarities.

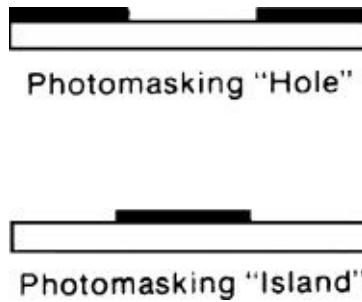


FIGURE 8.6 Photomasking hole and island.

The resist reaction to light just described is a character of negative-acting photoresists. There are also positive-acting photoresists. Within these resists, the light changes the chemical structure from relatively nonsoluble to much more soluble. The term describing this change is *photosolubilization*. [Figure 8.7](#) shows that an island is produced when a light-field mask is used with a positive photoresist.

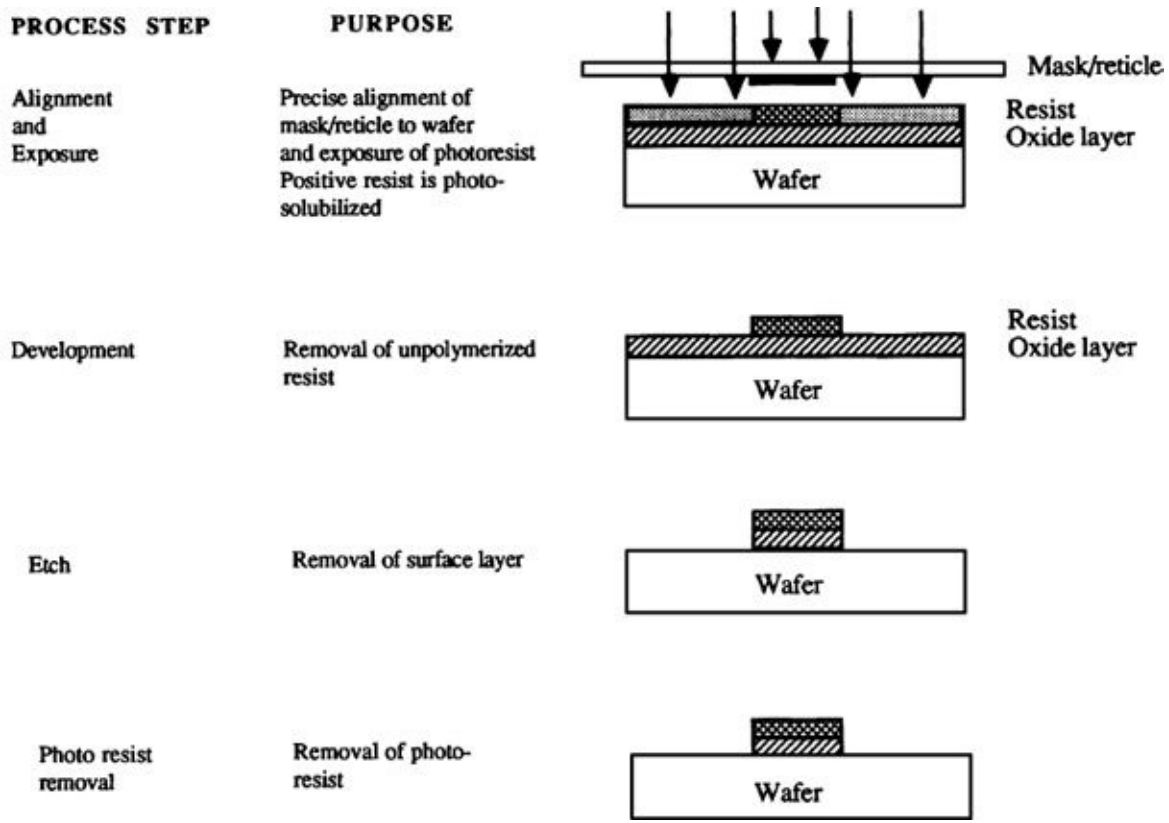


FIGURE 8.7 Image transfer from a light-field mask with a positive photoresist to create an island.

The result obtained from the photomasking process from different combinations of mask and resist polarities is shown in [Fig. 8.8](#). The choice of mask and resist polarity is a function of the level of dimensional control and defect protection required to make the circuit work. These issues are discussed in the remaining sections of the chapter.

		Photoresist Polarity	
		Negative	Positive
MASK POLARITY	Clear Field	HOLE	ISLAND
	Dark Field	ISLAND	HOLE

FIGURE 8.8 Mask and photoresist polarity results.

Ten-Step Process

Transferring the image from the reticle or mask onto the wafer surface layer is a

multistep procedure ([Fig. 8.9](#)). Feature size, alignment tolerance, the wafer surface, and the masking layer number all influence the difficulty and steps for a particular masking process. Many photo processes are customized to the particular conditions and outcomes. However, most are variations or options of a basic ten-step process. The process illustrated is shown with a light-field mask and a negative photoresist.

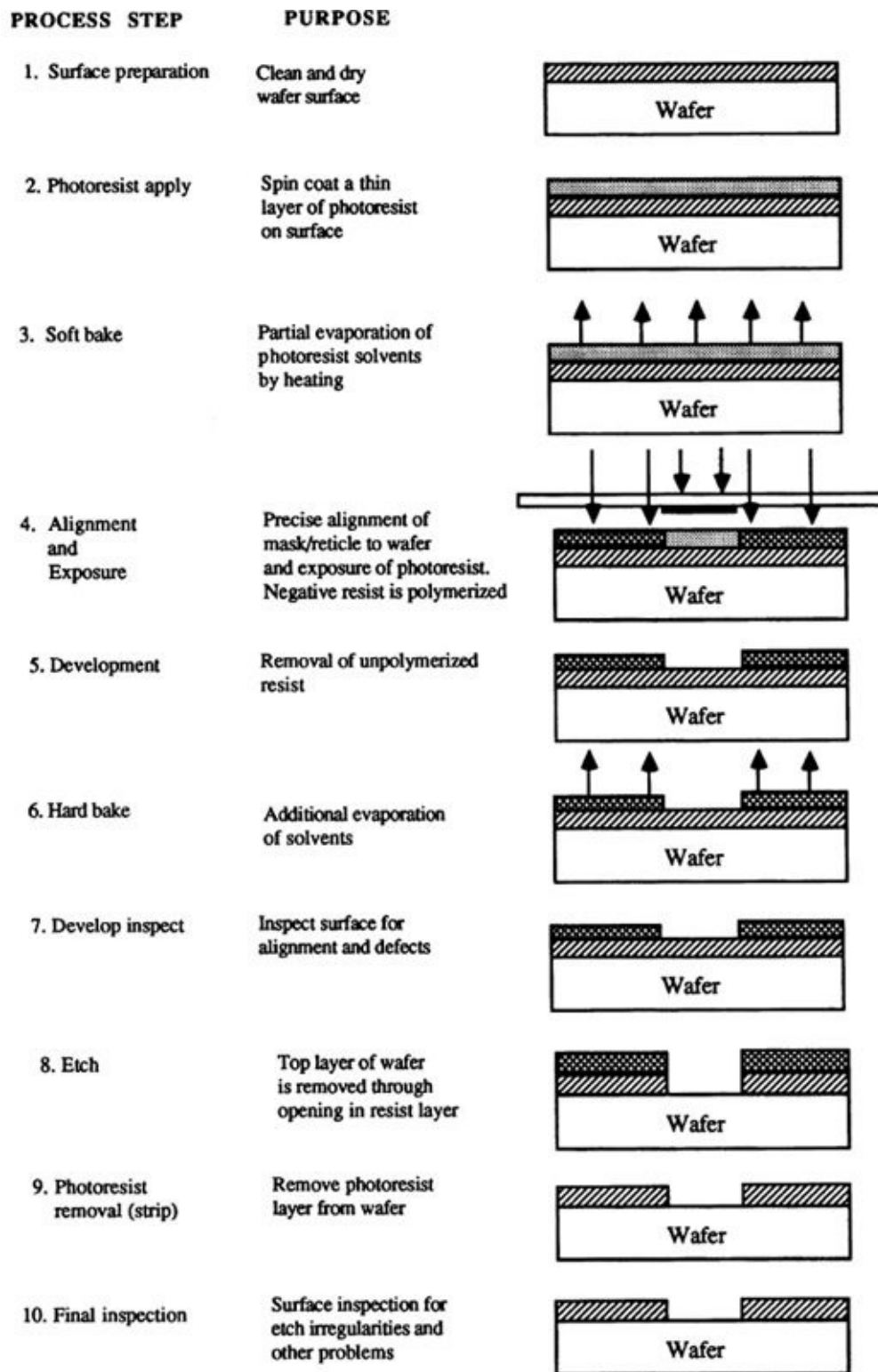


FIGURE 8.9 Ten-step photomasking process.

The first image transfer takes place in steps 1 through 7. In steps 8, 9, and 10, the image is transferred (second-image transfer) into the wafer surface layer. The reader is challenged to list the steps and draw the corresponding cross-sections using combinations of a dark-field mask and a positive photoresist. It is strongly

recommended that the reader master this ten-step process before proceeding to more advanced photolithography processes.

Basic Photoresist Chemistry

Photoresists have been used in the printing industry for over a century. In the 1920s, they found wide application in the printed circuit board industry. The semiconductor industry adapted this technology to wafer fabrication in the 1950s. Negative and positive photoresists designed for semiconductor use were introduced by Eastman Kodak and the Shipley Company, respectively, in the late 1950s.

The photoresist is the heart of the masking process. The preparation, bake, exposure, etch, and removal processes are fine-tuned to accommodate the particular resist used and the desired results. The selection of a resist and development of a resist process is a detailed and lengthy procedure. Once a resist process is established, it is changed very reluctantly.

Photoresist

Photoresists are manufactured for both general and specific applications. They are tuned to respond to specific wavelengths of light and different exposing sources. They are given specific thermal flow characteristics and formulated to adhere to specific surfaces. These properties come about from the type, quantity, and mixing procedures of the particular chemical components in the resist. There are four basic ingredients ([Fig. 8.10](#)) in photoresists: polymers, solvents, sensitizers, and additives (see [Chap. 10](#)).

Component	Function
Polymer	Polymer structure changes from soluble to polymerized (or vice versa) when exposed by the exposure source in an aligner.
Solvent	Thins resist to allow application of thin layers by spinning.
Sensitizers	Controls and/or modifies chemical reaction of resist during exposure.
Additives	Various added chemical to achieve process results, such as dyes.

FIGURE 8.10 Photoresist components.

Light-Sensitive and Energy-Sensitive Polymers

The ingredients that contribute the photosensitive properties to the photoresist are special light-and energy-sensitive polymers. Polymers are groups of large, heavy molecules containing carbon, hydrogen, and oxygen that are formed into a repeated pattern. Plastics are a form of polymers.

Resists designed to react to ultraviolet or laser sources are called *optical resists*. Other resists are designed to respond to X-rays or electron beams. In a negative resist, the polymers change from unpolymerized to polymerized after exposure to a

light or energy source. Physically, the polymers form a cross-linked material that is etch-resistant (see [Fig. 8.11](#)). Polymerization also happens when the resist is exposed to heat and/or normal light. To prevent accidental exposure, photomasking areas processing negative resist use yellow filters or yellow lighting.

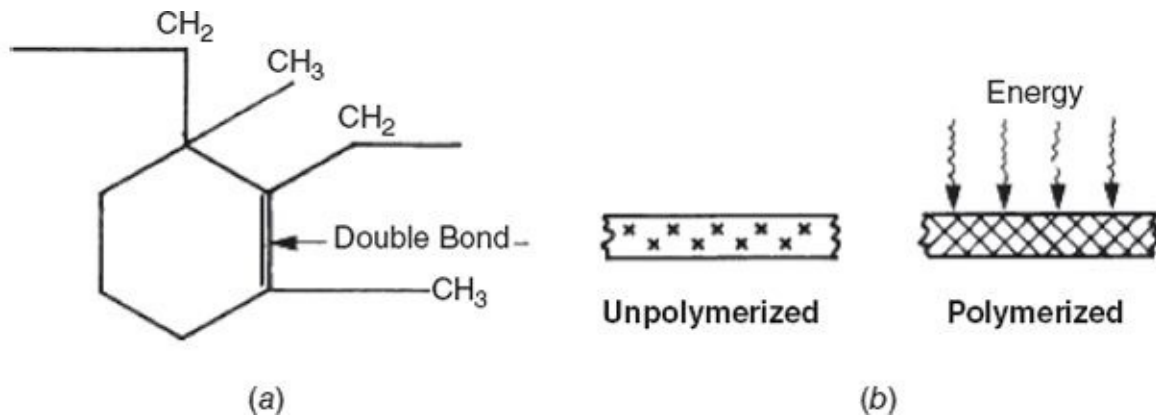


FIGURE 8.11 Negative resist chemistry.

The basic positive photoresist polymer is the phenol-formaldehyde polymer, also called the phenol-formaldehyde novolak resin ([Fig. 8.12](#)). Within the resist, the polymer is relatively insoluble. After exposure to the proper light energy, the resist converts to a more soluble state. This reaction is called *photosolubilization*. The photosolubilized part of the resist can be removed by a solvent in the development process.

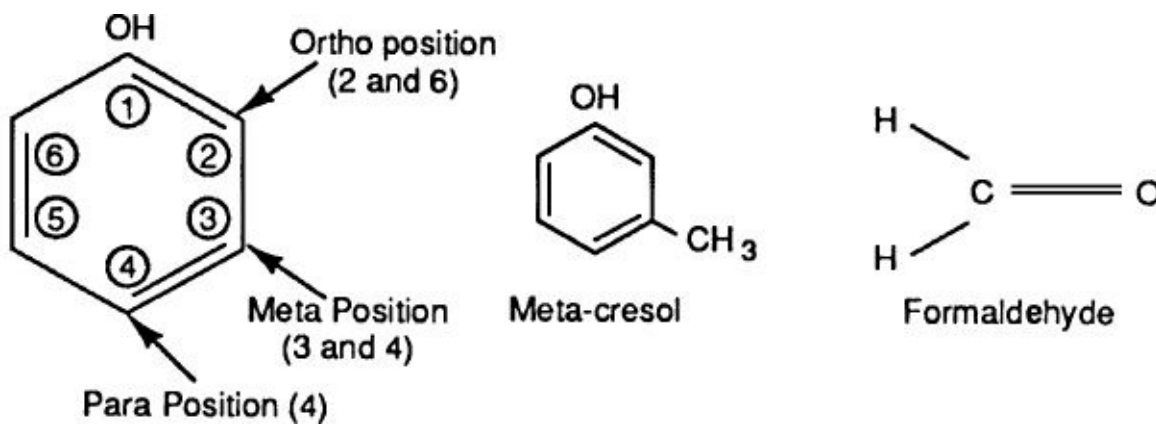


FIGURE 8.12 Phenol-formaldehyde novolak resin structure. (After: W. S. DeForest, *Photoresist: Materials and Processes*, McGraw-Hill, New York, 1975.)

Photoresists respond to many forms of energy. The forms are often referred to by their general category (light, heat radiation, and so on), or by a specific portion of the electromagnetic spectrum like ultraviolet light (UV), deep ultraviolet (DUV), I line, and so forth (see “Exposure Sources”). The exposing energies used are

detailed in the section on alignment and exposure. A number of strategies are used to resolve small images (see “Comparison of Positive and Negative Resists”). One is to use a narrower (or single) wavelength exposing source. The traditional novolak-based positive resist has been fine-tuned for use with I-line exposure sources. However, it does not work as well with DUV sources. Resist manufacturers have developed *chemically amplified resists* for this exposure source. Chemically amplified means that the chemical reactions of the polymer are increased by chemical additives. Resists for X-ray and electron beam (e-beam) are based on polymers different from conventional positive and negative resist chemistry.

Solvents

The largest ingredient by volume in a photoresist is the solvent. It is the solvent that makes the resist a liquid and allows the resist to be applied to the wafer surface as a thin layer by spinning. Photoresist is analogous to paint, which is composed of the coloring pigment dissolved in an appropriate solvent. For negative photoresist, the solvent is an aromatic type, xylene. In positive resist, the solvent is either ethoxyethyl acetate or 2-methoxyethyl.

Sensitizers

Sensitizers are added to either broaden the response range or narrow it to a specific wavelength. In negative resists, a compound called bis-aryldiazide is added to the polymer to provide the light sensitivity.¹ In positive resists, the sensitizer is o-naphthaquinonediazide.

Additives

Various additives are mixed with resists to achieve particular results. Some negative resists have dyes intended to absorb and control light rays in the resist film. Positive resists may have chemical *dissolution inhibitor systems*. These are additives that inhibit the dissolution of nonexposed portions of the resist during the development step.

Photoresist Performance Factors

The selection of a photoresist is a complicated procedure. The primary driving force is the dimensions required on the wafer surface. The resist must first have the capability of producing those dimensions. Beyond that, it must also function as an etch barrier during the etching step, a function that requires a certain thickness for mechanical strength. In the role of etch barrier, it must be free of pinholes, which also requires a certain thickness. In addition, it must adhere to the top wafer surface, or the etched pattern will be distorted, just as a paint stencil will give a sloppy image if it is not taped tightly to the surface. These, along with process latitude and step coverage capabilities, are resist performance factors. In the selection of a resist, the process engineer often must make tradeoff decisions between the various performance factors. The photoresist is one part of a complicated system of chemical processes and equipment that must work together to produce the image results and be productive, and the equipment must have an acceptable cost of ownership.

Resolution Capability

The smallest opening or space that can be produced in a particular photoresist is generally referred to as its *resolution capability*. On the wafer the most critical device or circuit dimension (CD) is the goal of a patterning process. The smaller the opening or space produced, the better the resolution capability. Resolution capability for a particular resist is referenced to a particular process, including the exposing source and developing process. Changing the other process parameters will alter the inherent resolution capability of the resist. Generally, smaller line openings are produced with thinner resist film thicknesses. However, a resist layer must be thick enough to function as an etch barrier and be pinhole-free. The selection of a resist thickness is a tradeoff between these two goals.

The capability of a particular resist relative to resolution and thickness is measured by its *aspect ratio* ([Fig. 8.13](#)). The aspect ratio is calculated as the ratio of the resist thickness to the image opening. As the industry requires smaller patterns, the factor of pattern density and shape become an influence on photoresist design. Small contact holes and high-density pattern areas, like in memory arrays, expose and develop differently as a result of reflection and chemical reaction factors. Consequently, there are available resists specifically designed for use in these situations.

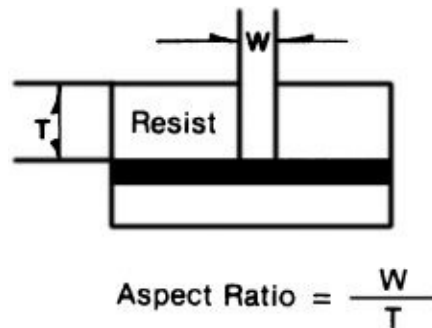


FIGURE 8.13 Aspect ratio.

Positive resists have a higher aspect ratio as compared to negative resists, which means that, for a given image-size opening, the resist layer can be thicker. The ability of positive resist to resolve a smaller opening is a result of the smaller size of the polymer. It is a little like using a smaller brush to paint a thinner line. It is the reason that positive resists are the choice for advanced high-density USLI level ICs.

Adhesion Capability

In its role as an etch barrier, a photoresist layer must adhere well to the surface layer to faithfully transfer the resist opening into the layer. Lack of adhesion results in distorted images. Resists differ in their ability to adhere to the various surfaces used in chip fabrication. Within the photomasking process, a number of steps are specifically included to promote the natural adhesion of the resist to the wafer surface. Negative resists generally have a higher adhesion capability than positive resists.

Photoresist Exposure Speed, Sensitivity, and Exposure Source

The primary action of a photoresist is a change in structure in response to an exposing light or radiation. An important process factor is the speed at which that reaction takes place. The faster the speed, the faster the wafers can be processed through the masking area.

The sensitivity of a resist relates to the amount of energy required to cause the polymerization or photosolubilization to occur. Furthermore, sensitivity relates to the energy associated with specific the wavelength of the exposing source. Understanding this property requires a familiarization with the nature of the electromagnetic spectrum ([Fig. 8.14](#)). Within nature, we identify a number of different types of energy: light, short and long radio waves, X-rays, and so on. In reality, they are all electromagnetic energy (or radiation) and are differentiated from each other by their wavelengths, with the shorter wavelength radiation having higher energies.

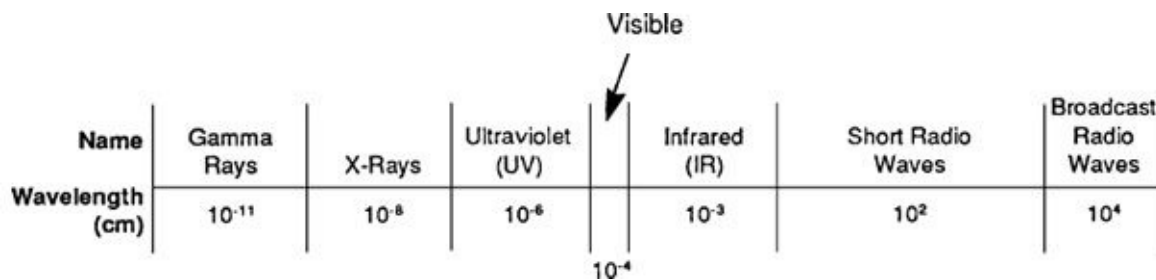


FIGURE 8.14 Electromagnetic spectrum.

Common positive and negative photoresists respond to energies in the ultraviolet and deep ultraviolet (DUV) portion of the spectrum ([Fig. 8.15](#)). Some are designed to respond to particular wavelength peaks (g, h, i lines in [Fig. 8.17](#)) within those ranges. Some resists are designed to work with X-rays or electron beams (e-beams). The industry has turned to lasers due to their high energy and narrow bandwidth.

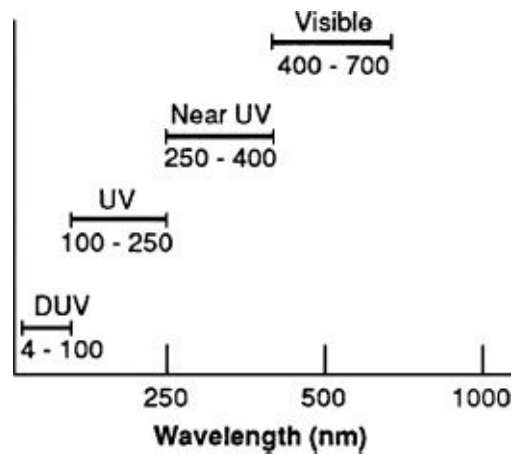


FIGURE 8.15 Ultraviolet and visible spectrum. (After Elliott.¹)

Resist sensitivity, as a parameter, is measured as the amount of energy required to initiate the basic reaction. The units are millijoules per square centimeter (mJ/cm^2).² The specific wavelengths the resist reacts to are called the *spectral response characteristics* of the resist. Figure 8.16 shows the spectral response characteristic of a typical production resist. The peaks in the spectrum are regions (wavelengths) that carry higher amounts of energy (Fig. 8.17). The different light sources used in masking areas are covered in the “Alignment and Exposure” section.

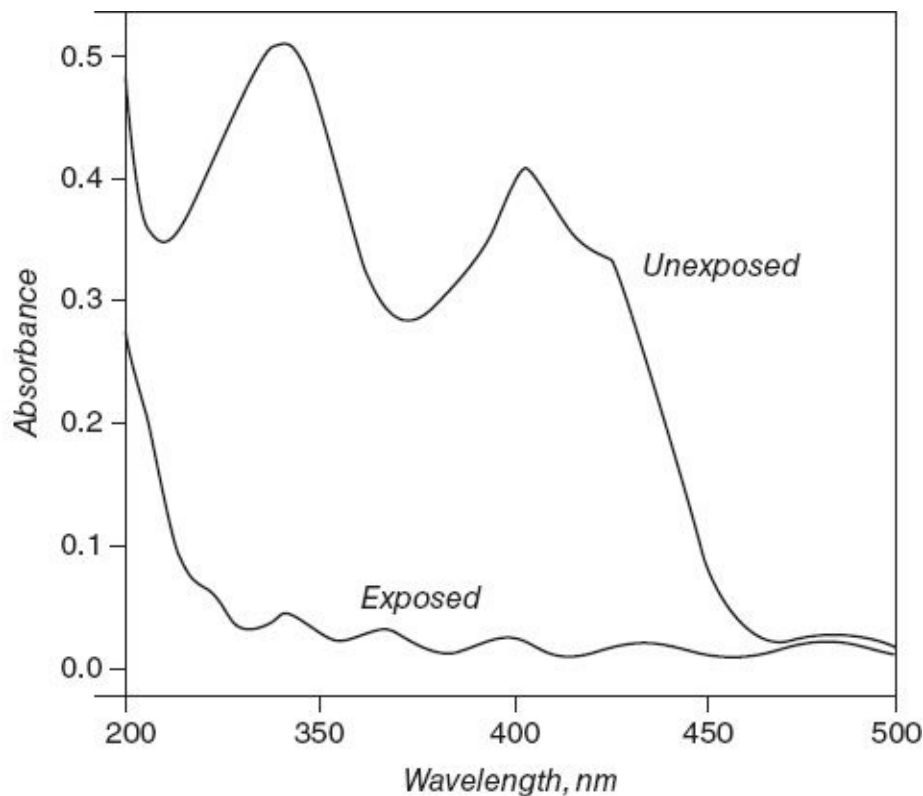


FIGURE 8.16 Absorbance curves.

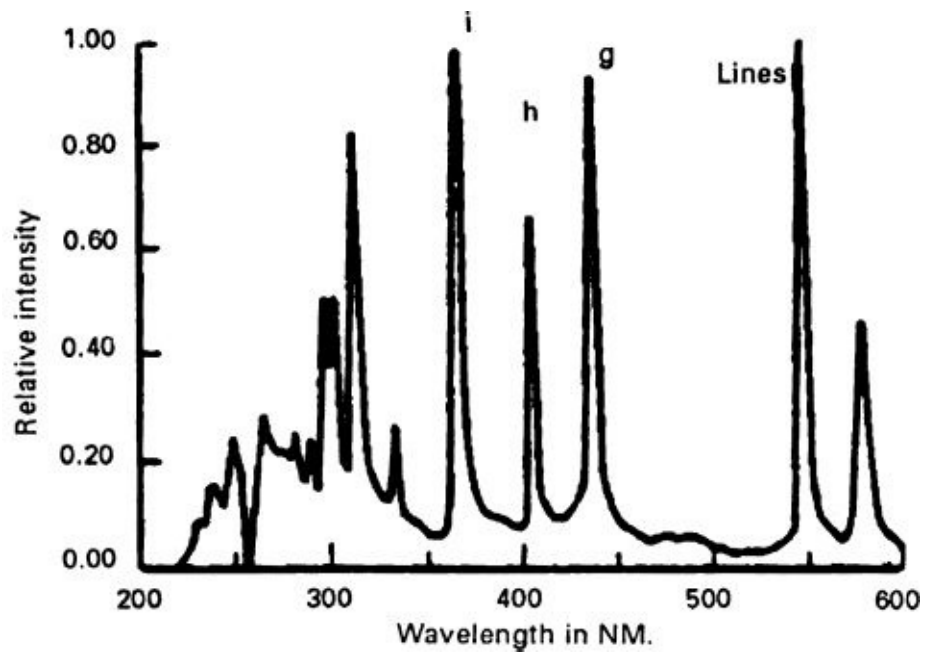


FIGURE 8.17 Mercury (Hg) spectrum. (From *Silicon Processing for the VLSI*, by Wolf and Tauber.)

Process Latitude

While reading the sections on the individual masking process steps, the reader should keep in mind the fact that the goal of the overall process is a faithful reproduction of the required image size in the wafer layer(s). Every step has an influence on the final image size, and each of the steps has inherent process variations. Some resists are more tolerant of these variations, that is, they have a wider process latitude. The wider the process latitude, the higher the probability that the images on the wafer will meet the required dimensional specifications.

Pinholes

Pinholes are microscopically small voids in the resist layer. They are detrimental, because they allow etchants to seep through the resist layer and etch small holes in the surface layer. Pinholes come from particulate contamination in the environment, the spin process, and from structural voids in the resist layer.

The thinner the resist layer, the more pinholes. Therefore, thicker films have fewer pinholes, but they also make the resolution of small openings more difficult. These two factors present one of the classic tradeoffs in determining a process resist thickness.

One of the principal advantages of positive resists is their higher aspect ratio, which allows a thicker resist film and a lower pinhole count for a given image size.

Particle and Contamination Levels

Resists, like other process chemicals, must meet stringent standards for particle content, sodium and trace metal contaminants, and water content.

Step Coverage

By the time the wafer is ready for the second masking process, the surface has a number of steps from added layers. As the wafer proceeds through the fabrication process, the surface gains more layers. For the resist to perform its etch barrier role, it must maintain an adequate thickness over these earlier layer steps. The ability of a resist to cover surface steps with adequate resist is also an important parameter.

Thermal Flow

During the masking process, there are two heating steps. The first, called *soft bake*, evaporates solvents from the resist. The second one, *hard bake*, takes place after the image has been developed in the resist layer. The purpose of the hard bake is to increase the adhesion of the resist to the wafer surface. However, the resist, being a plastic-like material, will soften and can flow during the hard-bake step. The amount of flow has an important effect on the final image size. The resist has to maintain its shape and structure during the bake, or the process design must account for dimensional changes due to thermal flow.

The goal is to achieve as high a bake temperature as possible to maximize adhesion without cause-resist flow that distorts the image. In general, the more stable the thermal flow of the resist, the better it is in the process.

Comparison of Positive and Negative Resists

Up to the mid-1970s, negative resist was dominant in the masking process. The advent of VLSI circuits and image sizes in the submicron range strained the resolution capability of negative resists. Positive resists had been around for over 20 years, but at that time their poorer adhesion properties were a drawback, and their superior resolution capability and pinhole protection were not needed.

In the 1980s, positive resist became the resist of choice. The transition was not easy. Switching to a positive resist requires changing the polarity of the masks or reticles. Reticle or mask dimensions print differently with the two resists ([Fig. 8.18](#)). With negative resist and light-field masks, the dimension in the resist is smaller than the mask/reticle dimension as a result of light wrapping (diffraction) around the image. With a positive resist and a dark-field mask, the diffraction tends to widen the image. These changes must be considered in making the mask/reticle and designing the other masking processes. In other words, an entirely new process is required to switch resist types.

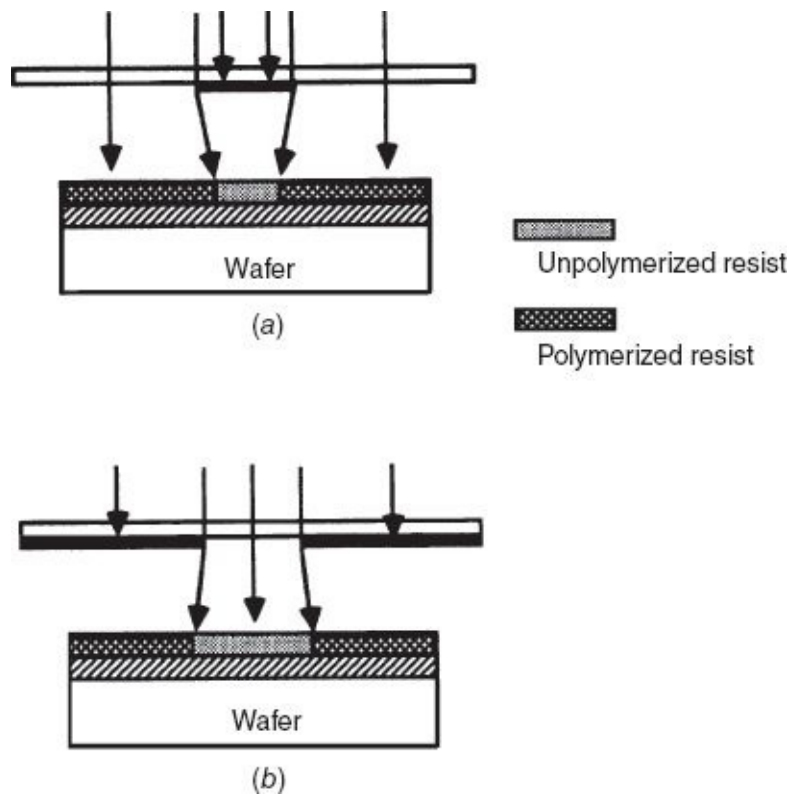


FIGURE 8.18 Changes in image size with (a) image size reduction with light-field mask and negative resist, and (b) image size enlargement with dark-field mask and positive photoresist.

Most of the images on most of the mask layers are holes. With positive resist, the mask polarity is dark field, which results in additional pinhole protection for the wafer ([Fig. 8.19](#)). Clear-field masks are prone to small cracks in the glass surface. These cracks, called *glass damage*, block the exposing light, creating unwanted holes in the photoresist layer, which in turn etch into the wafer surface as holes. The same is true for dirt particles that locate on the clear area of the mask/reticle. On dark-field masks, the majority of the surface is covered by chrome, which is hard and less likely to have pinholes. Thus, the wafer has fewer unwanted pinholes. For very small image sizes, positive resist is the only choice.

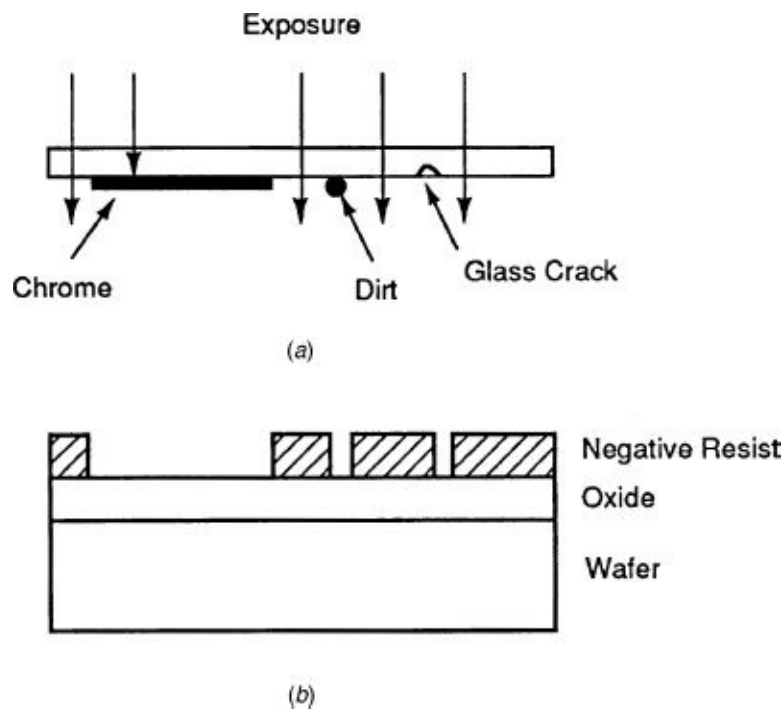


FIGURE 8.19 (a) Clear-field mask with dirt particle and glass crack, and (b) result in negative resist after develop.

Photoresist removal is also a factor. Generally, the removal of positive resists is easier and takes place in chemicals that are more environmentally sound. Fabrication areas producing simpler devices and circuits with larger image sizes or greater may still use negative resists. [Figure 8.20](#) shows a comparison of the properties of the two resists.

Parameter	Negative	Positive
Aspect Ratio (Resolution)		Higher
Adhesion	Better	
Exposure Speed	Faster	
Pinhole Count		Lower
Step Coverage		Better
Cost		Higher
Developers	Organic Solvents	Aqueous
Strippers		
Oxide Steps	Acid	Acid
Metal Steps	Chlorinated Solvent Compounds	Simple Solvents

FIGURE 8.20 Comparison of negative and positive results.

Physical Properties of Photoresists

The performance factors just detailed and all of the ten basic process steps are related to a number of physical and chemical properties of the resist. These properties are rigorously controlled by resist manufacturers.

Solids Content

A photoresist is a liquid that is applied to the wafer by a spinning technique. The thickness of resist left on the wafer is a function of the spin step parameters and several resist properties: *solids content* and *viscosity*.

Recall that the photoresist is a suspension of polymers, sensitizers, and additives in a solvent. Different resists will contain different amounts of these solids. The amount is referred to as the solids content of the resist and is expressed as the weight percent in the resist. Solids contents are in the 20 to 40 percent range.³

Viscosity

Viscosity is the quantitative measure of liquid flow. High-viscosity liquids, such as tractor oils, flow in a sluggish manner. Low-viscosity liquids, such as water, flow more readily. In both cases, the mechanism of flow is the same. The molecules in the liquid roll over each other as the liquid is being poured. As the molecules roll about, there is an attraction between them that acts as an internal friction. Viscosity is the measurement of that friction.

Viscosity is measured by several techniques. Most photoresist manufacturers measure viscosity with a rotating vane in the resist. The higher the viscosity, the more force required to move the vane through the liquid at a constant speed. The rotating-vane apparatus illustrates the force-related character of viscosity.

The unit of viscosity is the *centipoise* (one-hundredth of a poise). It is named after the French scientist Poiseuille, who investigated the viscous flow of liquids. One poise is equal to 1 dyne second per centimeter. The viscosity unit of centipoise is more correctly named the absolute viscosity.

Another unit, called the *kinematic viscosity*, is the centistoke. This value is calculated from the absolute viscosity (centipoise) divided by the density of the resist. Viscosity varies with temperature; therefore, its specified value is stated at a particular temperature, usually 25°C. Viscosity is a major parameter determining the resist thickness during the spin process. Viscosity is closely related to the solids content. The higher the solids content, the higher the viscosity.

Surface Tension

The surface tension of a resist also influences the outcome at spin. Surface tension is a measure of the attractive forces in the surface of the liquid ([Fig. 8.21](#)). Liquids with high surface tension flow less readily on a flat surface. It is the surface tension that draws a liquid into a spherical shape on a surface or in a tube.

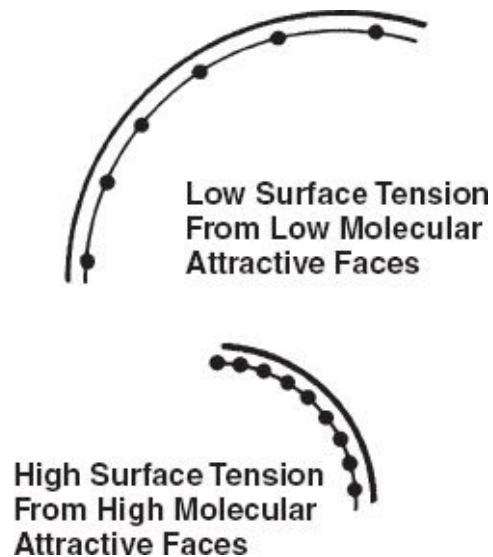


FIGURE 8.21 Surface tension.

Index of Refraction

The optical properties of the resist play a role in the exposure mechanism. One property is *refraction* or the bending of light as it passes through a transparent or semitransparent medium. Refraction is the same phenomenon that makes an object in a pool of water appear at a different location. This comes about as the light ray is slowed up by the material. The *index of refraction* is a measurement of the speed of light in a material compared to its speed in air, as shown in [Fig. 8.22](#). It is calculated as the ratio of the reflecting angle to the impinging angle. For photoresists, the index of refraction is close to that of glass, approximately 1.45.

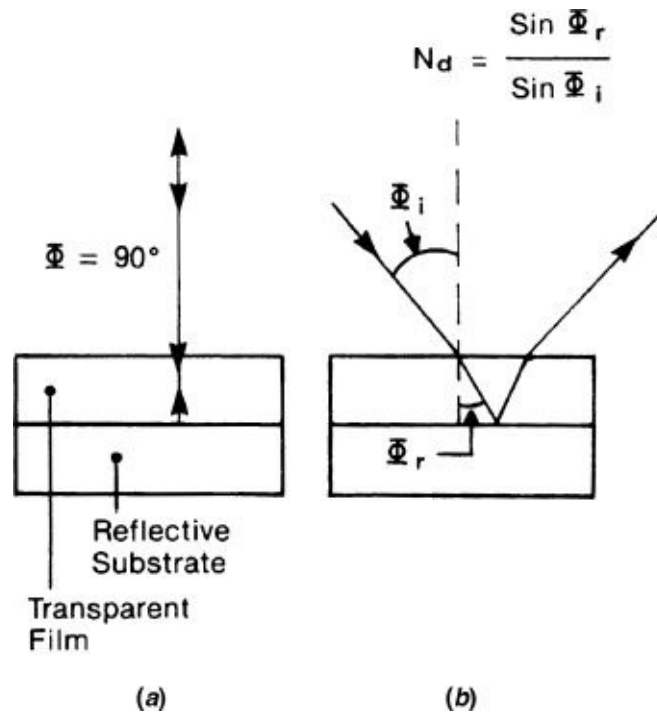


FIGURE 8.22 Index of refraction. (a) 90° incident light and (b) angled light refracted in the transparent film.

Storage and Control of Photoresists

Photoresists are delicate high-technology mixtures. Great care and precision go into their manufacture. Once a photomasking process is developed, its continuing success depends on the day-in, day-out control of the process parameters and a consistent photoresist product. Delivered batch-to-batch consistency is a responsibility of the manufacturer. Maintaining that consistency is the responsibility of the user. Several properties of resists dictate the required storage and control conditions.

Light and Heat Sensitivity

Both light and heat can activate the sensitive mechanisms in the resist. It is imperative that the resist be protected during storage and handling to prevent unwanted reactions that would interfere with the process results. This is the reason why masking areas have yellow lights when using negative resists. It is also the reason why resist bottles are brown. The colored glass protects the resist from stray light. Proper transportation and storage of resist require temperature control within the limits specified by the manufacturer.

Viscosity Sensitivity

Viscosity control is essential for good film thickness control. To maintain the photoresist's viscosity, resist containers must be sealed. Containers open to the atmosphere allow evaporation of the solvent, which results in a higher solid content, which in turn results in a higher viscosity. If photoresist is dispensed from plastic tubing, the material should be tested to ensure that the resist is not leaching plasticizers out of the material. The plasticizers will increase the viscosity of the resist.

Resist is also available in sealed vacuum pouches. The resist is protected during shipping and storage. During use, the pouch continues to collapse as the resist is dispensed, preventing air from reaching the resist surface.

Shelf Life

A container of photoresist comes with a recommended shelf life. The problem again has to do with the self-polymerization or photosolubilization of the resist. In time, changes to the polymer will take place, altering the resist performance when it reaches the production line.

Cleanliness

Needless to say, any and all equipment used to dispense photoresist must be maintained in the cleanest condition possible. Besides the effects of particular contamination from the system, the resist tubing must be cleaned regularly because of the possible buildup of dried photoresist. Cleaning agents should be checked for their compatibility with the resist. For example, trichloroethylene (TCE) should not be used with negative resist, because it can cause bubbles in the resist.

Photomasking Processes—Surface Preparation to Exposure

In the following sections, each of the ten basic masking steps is examined. Presented is the purpose of the step, technical considerations and challenges, options, and process control methods. Photomasking processes for advanced design rules use variations and different combinations of these basic process steps ([Chap. 10](#)).

Surface Preparation

Throughout this text, various analogies will be used to aid the reader in understanding complicated processes. A good analogy to the photomasking process is painting. Even the amateur painter soon learns that, to end up with a smooth film of paint that adheres well, the surface must be dry and clean. The same is true in photomasking technology. Ensuring that the resist will stick to the wafer surface requires surface preparation. This step is performed in three stages: particle removal, dehydration, and priming.

Particle Removal

The wafers coming to the photomasking area almost always come in a clean condition from another process such as oxidation, doping, or CVD. However, during storage, loading, and unloading into carriers, they may pick up some particulate contamination that must be removed from the surface. Depending on the level of contamination and/or the process requirements, several particulate removal techniques may be used ([Fig. 8.23](#)). In extreme cases, the wafers may be put through a wet chemical cleaning similar to the preoxidation cleaning processes, including an acid cleaning, water rinsing, and drying. The particular acid used must be compatible with the top layer on the wafer surface. Particle-removal methods are the same as described in [Chap. 7](#): manual blowoff, mechanical scrubbers, and/or high-pressure water spray.

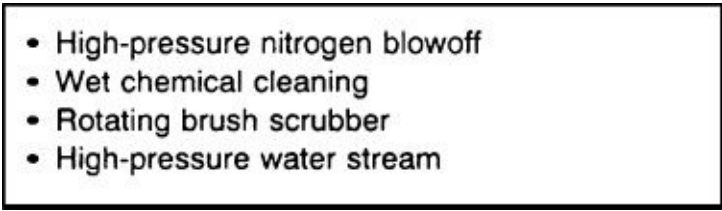
- 
- High-pressure nitrogen blowoff
 - Wet chemical cleaning
 - Rotating brush scrubber
 - High-pressure water stream

FIGURE 8.23 Prespin wafer-cleaning methods.

Dehydration Baking

It has been mentioned that the wafer surface has to be dry to promote adhesion. A dry surface is called *hydrophobic*, a chemical condition. Liquids form into small droplets on a hydrophobic surface, such as water beads on a newly waxed car. A hydrophobic surface is conducive to good photoresist adhesion ([Fig. 8.24](#)). Wafers coming to the masking operation usually have a hydrophobic surface.

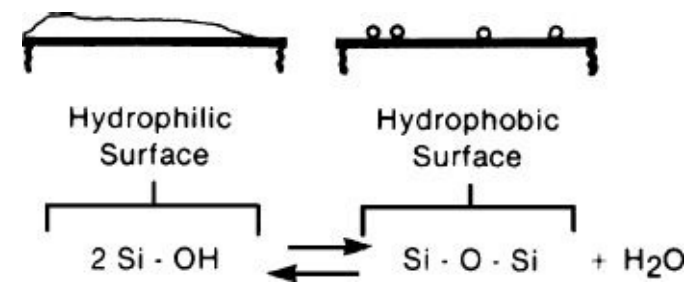


FIGURE 8.24 Hydrophilic versus hydrophobic surfaces.

Unfortunately, when the wafer is exposed to moisture, either from the air or from postcleaning rinses, the surface condition changes to a *hydrophilic* one. This condition is evidenced by a liquid on the surface spreading out in a wide puddle, such as water on a nonwaxed car surface. A hydrophilic surface is also said to be *hydrated*. Resist does not adhere very well to hydrated surfaces ([Fig. 8.24](#)).

Two important ways to maintain a hydrophobic surface are to keep the room humidity below 50 percent and to coat the wafers with photoresist as quickly as possible after being received from a previous process. Storage of the wafers in desiccators purged with dry, filtered nitrogen or in the dry mini-environment of a front opening unified pod (FOUP) box (see [Chap. 4](#)). Additional steps may be taken to establish a wafer surface with acceptable adhesion properties. These steps include a dehydration bake and priming with a chemical.

A heating operation may be used to reset the wafer surface to a dehydrated condition. Dehydration bakes take place in three temperature ranges to address three different dehydration mechanisms. In the range of 150 to 200°C (low temperature), surface water is evaporated. At 400°C (medium temperature), water molecules loosely attached to the surface will leave. At temperatures above 750°C (high temperature), the surface is chemically restored to a hydrated condition.

In most masking processes, only low-temperature dehydration bake temperatures are used. This is because this temperature range is easily obtainable with track hot plates and chest-type convection and vacuum ovens. Another advantage of low-temperature dehydration is that the wafers do not have to wait for a cooldown

before the spin process. Systems to perform this step can easily be integrated into a spin-bake system, making them dehydration-spin-bake systems. An explanation of these heating systems is provided in the “Soft Bake” section.

High-temperature dehydration bakes are rare. One reason is that reaching a temperature of 750°C usually requires the use of a tube furnace or rapid thermal process unit that complicates the production flow. A second reason is the temperature level itself. At 750°C, doped junctions in the wafer can move (which is undesirable), and mobile ionic contaminants on the surface can move into the wafer causing reliability and performance problems.

Wafer Priming

In addition to dehydration baking, the wafers may go through a chemical priming step to ensure good adhesion of the resist. In painting, primers are a subcoat selected for their ability to adhere to the surface and provide a good surface to which the paint will stick. In semiconductor photomasking, the primer effect is similar. The primer chemically ties up molecular water on the wafer surface, thereby increasing its adhesion property.⁴

A number of chemicals provide priming capabilities, but hexamethyldisilazane (HMDS) has become the standard. The use of HMDS is described in U.S. Patent 3,549,368 by R. H. Collins and F. T. Deverse (1970) of IBM. The HMDS is mixed with xylene in solutions from 10 to 100 percent. The exact mixture is determined by the particular surfaces and environmental factors in the cleanroom. Unlike a painting primer, a thickness of only several molecules is sufficient to provide the necessary adhesion promotion.

Spin Priming

In most photomasking areas, liquid primers are applied to the wafer while it is on the resist spinner chuck (Fig. 8.25). Automatic spinners (see next section) have a separate system to dispense HMDS onto the wafer surface just prior to the application of the resist. After the spinner dispenses the primer onto the rotating wafer, the chuck is ramped to a higher speed to dry the HMDS layer. A major production advantage of spin priming is that it takes place in-line with the spin step.

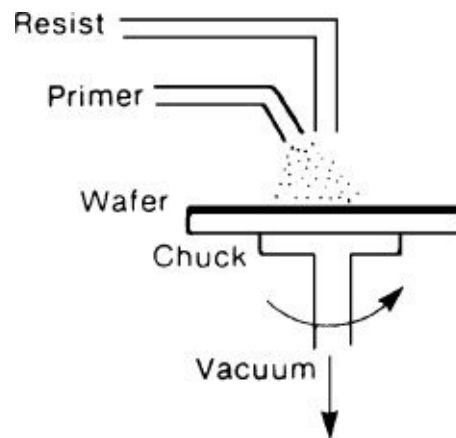


FIGURE 8.25 Spin dispense of primer.

Vapor Priming

Both immersion and spin priming require the direct contact of the HMDS liquid with the wafer surface. Whenever a liquid is in contact with the wafer, there is a danger of contamination from the liquid. Another consideration is that the HMDS must be dry before the resist is applied. Wet HMDS can dissolve the bottom layer of the resist and interfere with the exposure, development, and etching. Lastly, HMDS is relatively expensive, and in spin priming, an excess of HMDS is sprayed on the wafer to ensure adequate coverage. The excess is thrown off the wafer and discarded.

The preceding considerations are overcome by vapor priming techniques. Vapor priming is practiced in three forms. Two are performed at atmospheric pressure and one in a vacuum ([Fig. 8.26](#)). One atmospheric system employs a bubbler chamber connected to a desiccator-type chamber. Nitrogen is bubbled through the HMDS and carried into the chamber, where it coats the wafers. Another method employs a vapor degreaser with a reserve of liquid HMDS. The HMDS is heated to the vapor point and the wafers are suspended in the vapors for coating.

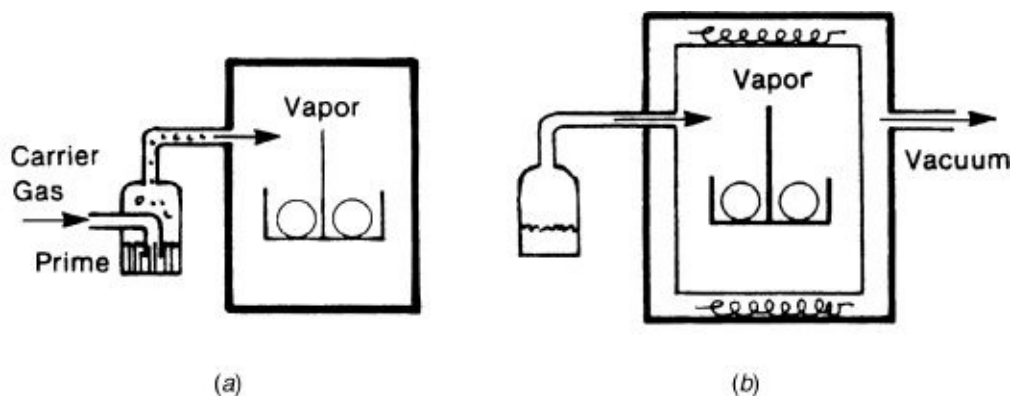


FIGURE 8.26 Vapor prime methods: (a) atmospheric and (b) vacuum bake-vapor prime.

The third vapor technique is vacuum vapor priming, which uses a sealed flask of HMDS connected to a vacuum oven or single-wafer chamber. The wafers are first heated in the oven in a nitrogen atmosphere. After a temperature of about 150°C is reached, the atmosphere is switched to a vacuum. Once the vacuum level is reached, a valve is opened, and HMDS vapors are drawn into the chamber by the low pressure. Within the chamber, the wafers become completely coated as the vapors fill the entire chamber. This method has shown good adhesion longevity, even in the presence of high humidity.

Vacuum vapor priming offers the additional advantage of a combined dehydration bake and prime step and a significant reduction in HMDS usage.

Vacuum vapor priming practiced in a chest-type oven adds an additional step to the process. Many automatic spinner systems incorporate in-line vacuum vapor primer units.

Photoresist Application (Spinning)

The purpose of the photoresist application step is the establishment of a thin, uniform, defect-free film of photoresist on the wafer surface.

These goals are easy to state, but they require sophisticated equipment and stringent controls to achieve. A typical resist layer varies from 0.5 to 1.5 μm in thickness and has to have a uniformity of only $\pm 0.01 \mu\text{m}$ (100 Å).

The usual methods of applying thin layers of liquids to surfaces are brushing, rolling, and dipping. None of these methods is adequate to achieve the quality resist film necessary for photomasking. The method used is spinning, which was briefly described in the section on priming. Spinners are built in manual, semiautomatic, and automatic designs. The systems differ in the degree of automation and are described in the following text. However, the deposit of the film on the wafer is common to each of the systems.

Spin processes are designed to prevent or minimize the buildup of a bead of resist around the outer edge of the wafer. Called an *edge bead*, the buildup causes image distortions during the exposure and etch processes.

The Static Dispense Spin Process

Application of the resist occurs immediately after the priming process. The wafer is placed on a vacuum chuck and several cubic centimeters (cm^3) of the photoresist is deposited in the center of the wafer ([Fig. 8.27](#)) and allowed to spread out into a puddle. The puddle is allowed to spread until it covers the majority of the wafer surface. The amount of resist deposited in the puddle is critical only in the extremes. Too small an amount will result in incomplete resist coverage, and too much will cause a buildup of a resist rim or result in resist on the back of the wafer ([Fig. 8.28](#)).

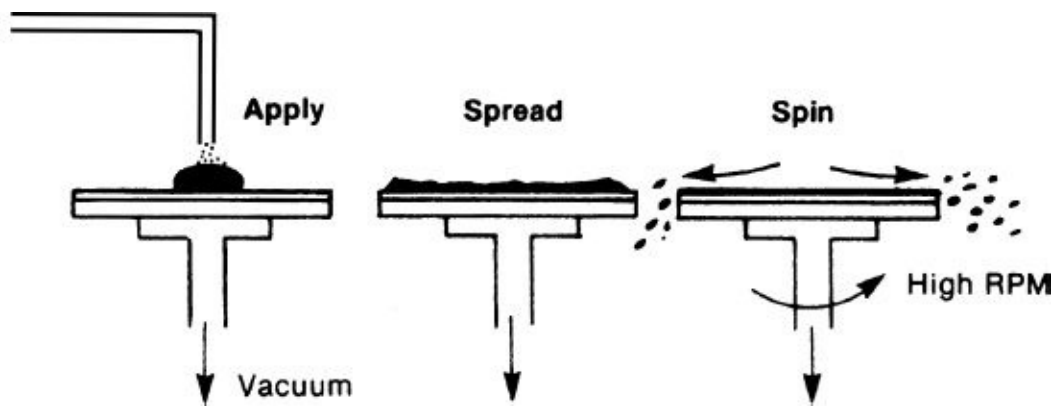


FIGURE 8.27 Static spin process.

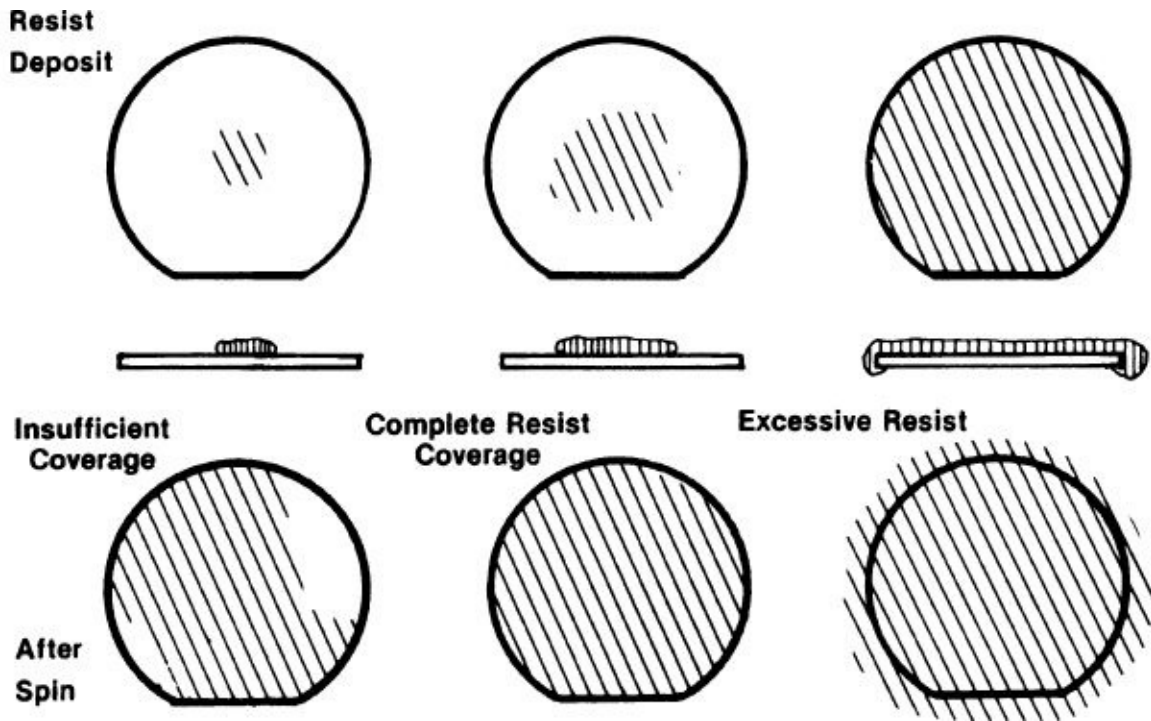


FIGURE 8.28 Example of resist coverage.

When the puddle reaches its specified diameter, the chuck is rapidly accelerated to a predetermined speed. During the acceleration, centrifugal forces spread the resist to the wafer edge and throw off excess resist, leaving a thin uniform layer on the wafer. The high-speed spin continues for some time after the resist is spread to allow drying of the resist.

The final thickness of the film is established as the result of the resist viscosity, the spin speed, the surface tension, and the drying characteristics of the resist. In practice, surface tension and the drying characteristics are properties of the resist, and the viscosity-spin speed relationship is determined from curves supplied by the resist manufacturer or established for the particular spin system used. [Figure 8.29](#) shows a typical thickness versus spin speed relationship.

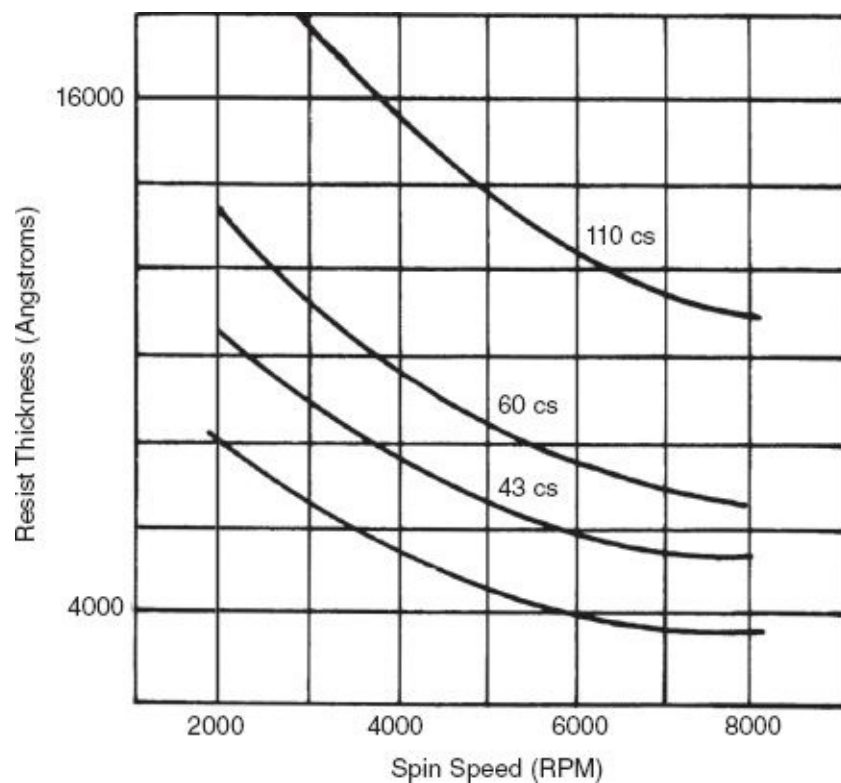


FIGURE 8.29 Typical resist thickness versus spin speed.

Although spin speed is specified to control resist thickness, it is actually the acceleration rate that establishes the final resist thickness. The acceleration characteristic of the spinner must be specified and maintained as a constant in the spin process.

Dynamic Dispense

The need for uniform resist films on larger-diameter wafers led to the development of the dynamic spin dispensing technique ([Fig. 8.30](#)). For this technique, the wafer is rotated at a low speed of approximately 500 rpm. While the wafer is rotating, the resist is dispensed onto the surface. The action of the rotation assists in the initial spreading of the resist. Less resist is used, and a more uniform layer is achieved. After spreading of the resist, the spinner is accelerated to a high speed to complete the spread and thin the resist into a uniform film.

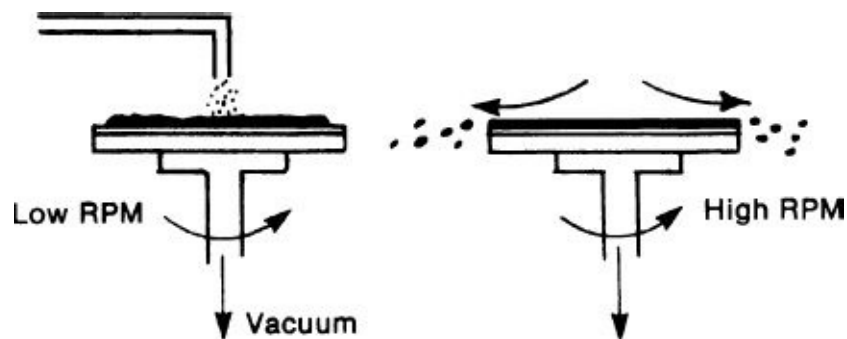


FIGURE 8.30 Dynamic spin dispense.

Moving-Arm Dispensing

An improvement on the dynamic dispense technique is the addition of a moving-arm resist dispenser ([Fig. 8.31](#)). The arm moves in a slow motion from the center of the wafer toward its edge. This action creates more uniform layers. A moving-arm dispenser also saves resist material, especially for larger-diameter wafers.

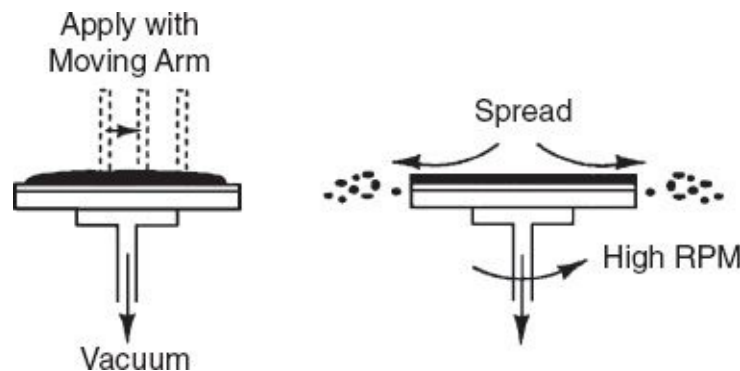


FIGURE 8.31 Moving-arm dispensing.

Manual Spinners

A manual spinner is a simple machine consisting of from one to four vacuum chucks (called *heads*), a motor, a tachometer, and a connection for a vacuum source. They are used in small labs and teaching facilities. Each head is surrounded by a catch cup, which serves to collect the excess resist and direct it to a collection vessel. The catch cup also prevents “balls” of resist thrown off the wafers during the acceleration from landing on adjacent wafers. The process starts with the removal of particles from the wafer with a filtered nitrogen gun. The wafer(s) are mounted on the head with tweezers or a vacuum wand, and the chuck vacuum is turned on. Next, the HMDS is dispensed onto the surface from a syringe or squeeze bottle, and the wafer is spun and dried. Finally, the resist puddle is dispensed from another syringe or squeeze bottle. Most spin processes performed on manual spinners use the static dispense method.

Automatic Spinners

A semiautomatic spinner adds automation to the wafer blowoff, resist dispense, and spinning cycles. The nitrogen blowoff is accomplished from a separate tube over the vacuum chuck that is connected to a pressurized nitrogen source ([Fig. 8.32](#)). Also in the dispense chamber are a primer dispense tube and a resist dispense tube. The resist tube is fed resist from either a nitrogen pressurized vessel or through a diaphragm-type pump. In general, the industry has moved away from pressurized, resist dispense systems because of problems that arise from the absorption of nitrogen into the resist. The nitrogen comes out of the resist after the dispense cycle, causing voids in the film. Diaphragm pumps eliminate this problem.

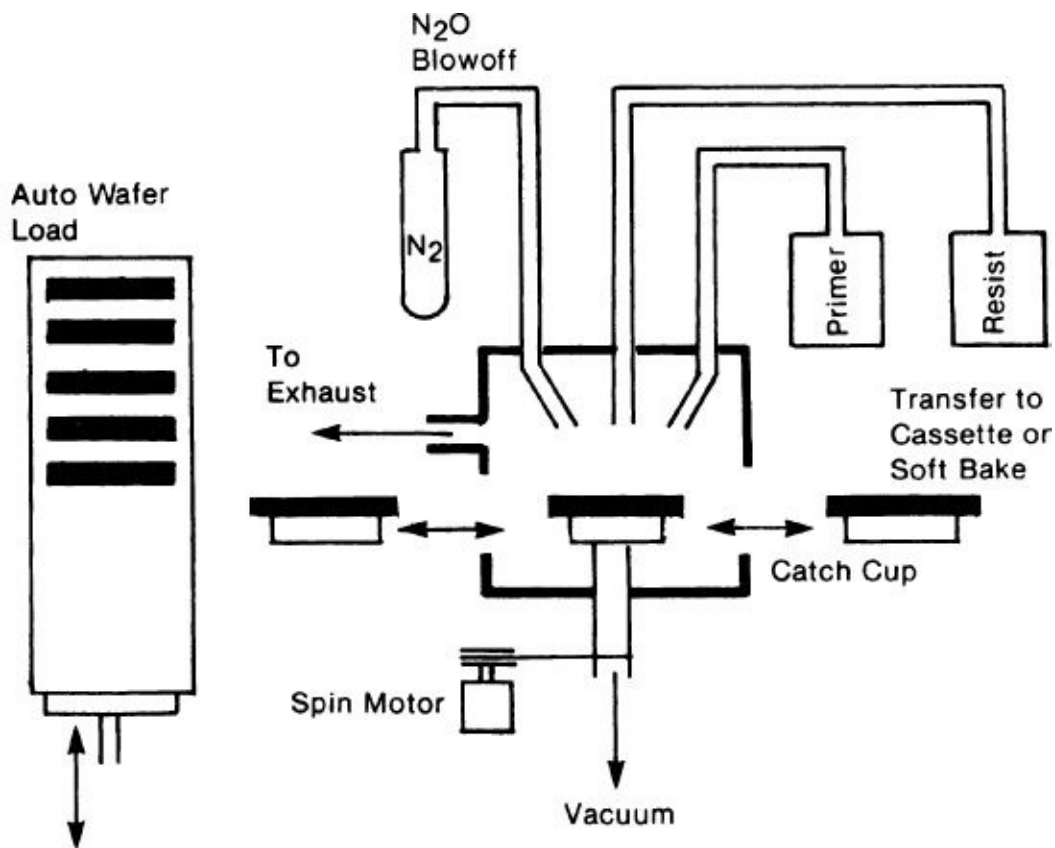


FIGURE 8.32 Automated spinner diagram.

Automatic resist dispensers have a negative-pressure capability that automatically draws the resist back up into the dispenser tube after each dispensing operation. This *drawback* (or *suckback*) minimizes the exposed surface of the resist in the tube from drying into a hard ball that can be deposited on the wafer. In fully automatic systems, all of the events of the spin process are controlled by microprocessors. The systems

have mechanisms to extract the wafers from the carriers, place them on the chucks, perform the priming step, dispense the photoresist, remove the edge bead, perform a soft bake, and place the wafers back in their carriers. The standard system configuration is in a line, referred to as a *track*. Production-level spinners will have from two to four side-by-side tracks.

Edge Bead Removal

A consequence of the high-speed spin can be a buildup of resist around the edge of the wafer, called an *edge bead*. It is removed by a solvent spray directed on the front and back of the wafer, near the edge ([Fig. 8.33](#)).

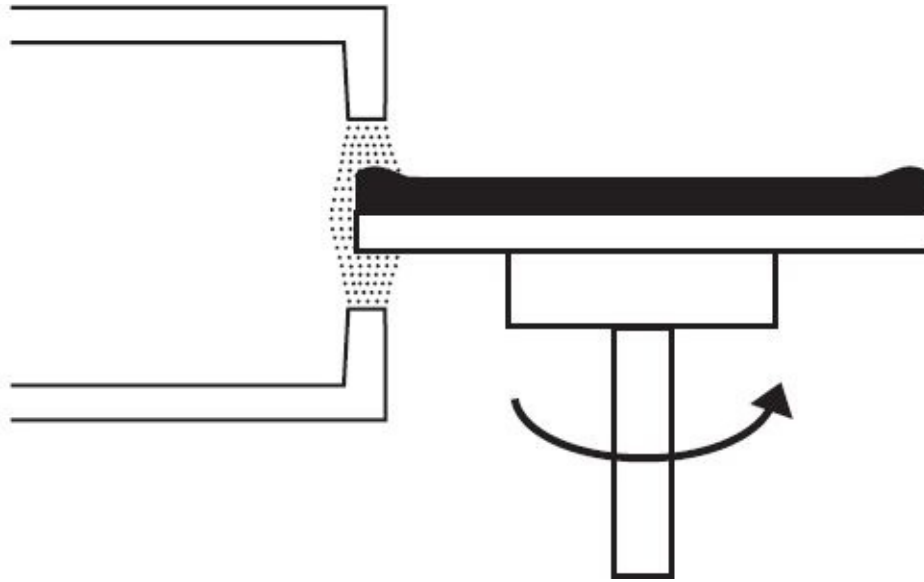


FIGURE 8.33 Chemical edge bead removal.

Backside Coating

In some device processes, it is required that the oxide on the back of the wafer remain in place through the masking process. One way this is accomplished is by coating the back of the wafer with photoresist. The requirement for this backside coating is simply a thick enough layer to survive the etch process.

Soft Bake

The soft-bake step is a heating operation with the purpose of evaporating a portion of the solvents in the photoresist. After the bake, the resist film is still “soft,” as opposed to being baked to a varnish-like finish. The solvents are evaporated for two reasons. Remember that the principal role of the solvents is to allow the application of the resist in a thin layer on the wafer. After that role is fulfilled, the presence of the solvent can interfere with the rest of the processing, particularly the exposure step. The solvent(s) in the resist can absorb exposing radiation, thus interfering with the proper chemical change in the photosensitive polymers. The second problem is with the resist adhesion. Using the painting analogy, we know that complete drying (evaporation of the solvent) is necessary for good adhesion. Photoresist adhesion is similar.

The principle soft-bake parameters are time and temperature. Two major goals of the patterning process are correct image definition and the adhesion of the resist to the wafer during the etch step. Both of these goals are influenced by soft bake. In the extreme, an underbaking will result in incomplete image formation at exposure and excessive lifting (poor adhesion) at the etch step. Overbaking will cause the polymers in the resist to polymerize and not react to exposing radiation.

Temperature and time ranges for the soft bake are provided by the resist manufacturer and are fine-tuned by the masking engineer. Positive photoresists can be bakes in air where negative photoresists have to be baked in a nitrogen atmosphere. A number of different methods are used to accomplish the soft-bake step, and various pieces of equipment incorporating all three heat-transfer methods are used.

The three methods of heat transfer are conduction, convection, and radiation. Conduction is the transfer of heat by direct physical contact of the object with a heated surface. Hot plates heat by conduction (labs and teaching facilities). In the conduction process, the vibrating atoms of the hotter surface cause the object atoms to vibrate also. As they vibrate and collide, the atoms become hotter.

Some dehydration baking is in convection ovens. Systems using convection heating include home forced-air furnaces, hair dryers, air-and nitrogen-fed ovens, and oxidation furnaces. In these systems, a unit heats the gas and a blower or pressure pushes the gas to a space where it, in turn, heats the object.

The third method is radiation. The term *radiation* describes the travel of electromagnetic energy waves through space. Radiation waves travel in vacuums as well as through gases. The sun transfers heat to the earth by radiation. Heating lamps also transfer heat by radiation. Radiation is the heating method used in rapid thermal processing (RTP) systems. When the radiation strikes an object, the energy

carried by the wave is transferred directly to the atoms of the object.

Convection Ovens

The mainstay baking oven for nonautomated fabrication lines is the convection oven ([Fig. 8.34](#)). It is a stainless-steel chamber in an insulated enclosure. Either nitrogen or air is supplied by ducts surrounding the chamber and passed through a heater before being directed into the chamber. The inside of the chamber is fitted with racks for the wafer carriers. The carriers stay inside the oven for a predetermined time while the heated gas brings them up to temperature. Convection ovens used for VLSI applications have proportional band controllers and HEPA filters to maintain a clean baking environment.

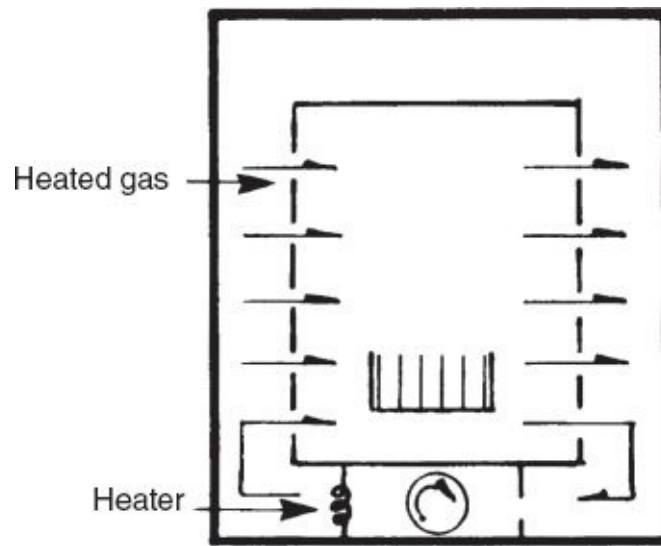


FIGURE 8.34 Convection oven.

There are several drawbacks to convection ovens for soft baking. One is batch-to-batch temperature variation, which arises from the amount of time the door is open for loading, the size of the load, and the variable time for all parts of the oven to reach a constant temperature. A process problem associated with convection heating is the tendency of the top layer of resist to “crust,” trapping solvents in the resist ([Fig. 8.35](#)).

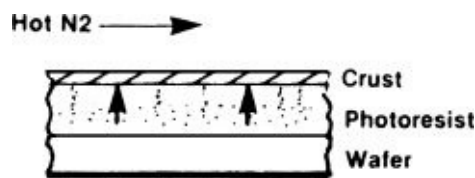


FIGURE 8.35 Crusting effect of ovens.

Manual Hot Plates

In manual and laboratory operations, a simple hot plate can be used for the soft bake. The wafers are placed on an aluminum holder, which is set on the hot plate ([Fig. 8.36](#)). A process advantage of hot-plate heating is that the bottom of the wafer is heated first. This allows the solvents to escape through the top surface and also minimizes “crusting” of the surface. A hot-plate process is operator-dependent and has low productivity.

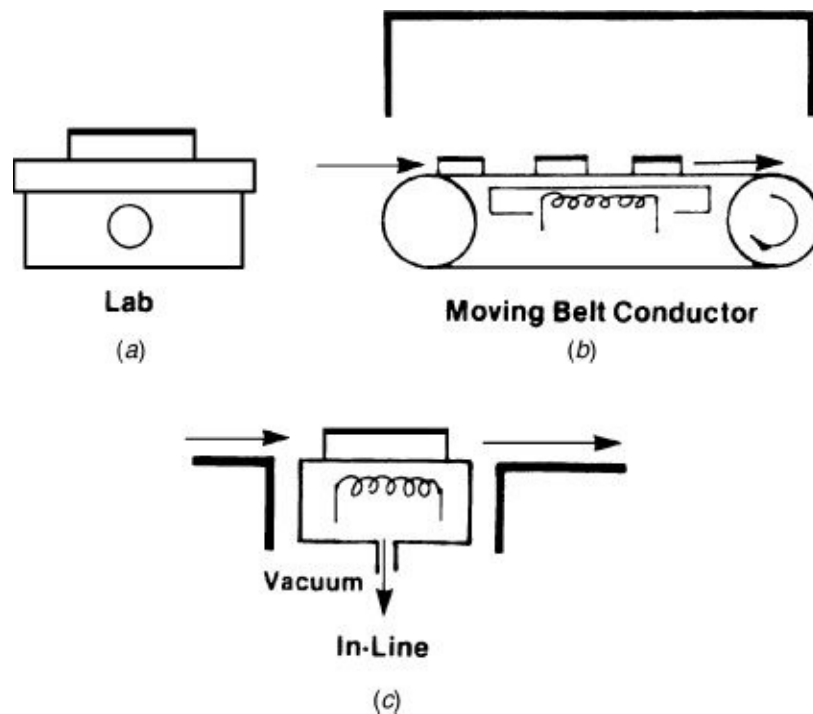


FIGURE 8.36 (a) Manual hot plate, (b) in-line continuous hot plate, and (c) in-line single-wafer hot plate.

In-Line, Single-Wafer Hot Plates

The backside advantage of hot-plate heating can be gained in track systems when a single-wafer hot plate is built into the system. Wafers leaving the spinner are positioned on the hot plate and clamped to it with a vacuum. The wafer and resist are heated for a predetermined time, the vacuum released, and the wafer transferred to a carrier. These in-line systems are connected to the facility exhaust system for the removal of the solvent vapors.

Moving-Belt Hot Plates

A constraint on the productivity of single-wafer hot plates in an integrated system is

the total time of the spin step. A typical spin time is 25 to 40 s, which means the soft bake would have to be completed in that amount of time to keep the wafers moving in a continuous flow. For some resists and for some processes, that is too short a time. A way around the problem is a moving-belt hot plate. The wafers are placed on a heated, moving steel belt, and the temperature and belt speed are set to meet the soft-bake requirements and process the wafers in a continuous flow.

Moving-Belt Infrared Ovens

The desire for fast, uniform, and noncrusting soft-bake methods led to the development of infrared (IR) radiation sources ([Fig. 8.37](#)). Infrared baking is much faster than conduction baking and heats from the “inside out.” Inside-out baking is the principle of conduction hot-plate baking. The infrared waves pass through the resist layer without heating it, much like sunlight will pass through a window without heating it. The wafer, however, absorbs the energy, gets hot, and in turn heats the resist layer from the bottom.

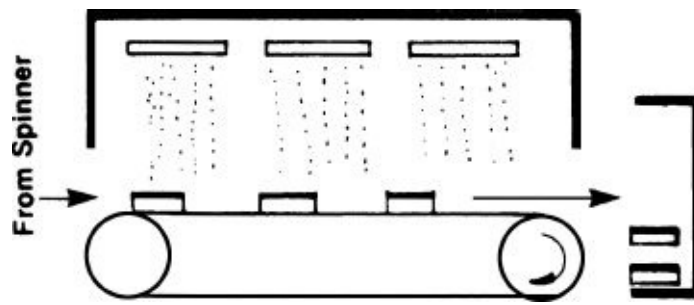


FIGURE 8.37 Moving-belt infrared (IR) heating.

Microwave Baking

Microwaves as a soft-bake heating source have the advantage of infrared heating but at a much faster rate because of the higher energy carried in a microwave. Soft-bake temperatures can be well under 1 min. This brief time lends itself to on-chuck soft baking. Immediately after spinning, a microwave source is directed at the wafer, completing the soft bake ([Fig. 8.38](#)).

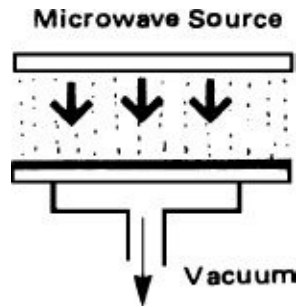


FIGURE 8.38 Microwave heating.

Vacuum Baking

Vacuum ovens offer several advantages for a number of process steps. A vacuum oven is configured similarly to a convection oven, but is fitted to a vacuum source. Vacuum is particularly efficient for evaporation processes, because the reduced pressure aids the evaporation of the solvents, reducing the reliance on the temperature. However, for soft baking, the wafers must be heated to a uniform temperature. A problem arises because heating in a vacuum oven is by radiation from the heated chamber walls to the wafers. This heat transfer method is sometimes called *line-of-sight* because, for uniformity, each wafer must have a clear line of sight to the heat source. In a chest-type vacuum oven packed with carriers of vertically held wafers, this condition cannot be met. The result is poor temperature uniformity in most vacuum ovens.

The benefits of vacuum and uniform hot-plate heating can be achieved with in-line, single-wafer systems. [Figure 8.39](#) is a table summarizing the different soft-bake methods.

Method	Bake Time Min.	Temperature Control	Productivity Type
Hot plate	5–15	Good	Single to small batch
Convection oven	30	Average - good	Batch
Vacuum oven	30	Poor average	Batch
I.R. moving belt	5–7	Poor average	Single
Conductive moving belt	5–7	Average	Single
Microwave	0.25	Poor average	Single

FIGURE 8.39 Soft-bake chart.

Post-Soft-Bake Cooling (Chill)

Track systems will often move the wafers to a cooling or chilling station. The cooldown is temperature controlled. The chill plates feature solid-state cooling technology.

Alignment and Exposure

Alignment and exposure (A&E) is one of the basic ten-patterning steps, with two separate actions. The first part of the A&E step is the positioning or alignment of the required image on the correct location on the wafer surface. The second part is the encoding of the image in the photoresist layer from an exposing light or other radiation source. Correct alignment of the image patterns and establishment of the precise image dimensions in the resist are absolute requirements for the functioning of the devices and circuits. Furthermore, the wafers spend 60 percent of the process time in the lithography area.

Besides image shape and dimensional control pattern overlay or alignment tolerance is a critical parameter. Again the skyscraper analogy is useful. If the elevator shafts from floor to floor do not align there will be serious problems with constructing a functioning elevator system. In integrated circuits the same principle applies. Individual patterned layers must fit inside and envelop to ensure device functioning. The exact tolerance is a percentage of the critical dimensions and pitch of the particular circuit. And the overlay tolerance is challenged and becomes more critical as the number of patterning layers climbs. Various formulas are used as discussed in the 2011 version of the ITRS section on Lithography.²

Alignment and Exposure Systems

Up to the mid-1970s, the photoresist engineer had the choice of only two A&E systems: contact and proximity aligners. Both of these use the mask-patterning system. The required layer pattern is first created in a chrome layer on a glass substrate. This pattern is in turn transferred into the photoresist layer. A light source shines through the mask and codes the pattern as polymerized and unpolymerized areas in the photoresist layer.

Alignment with a mask system requires positioning the mask either to the blank wafer (first mask) or to the patterns already on the wafer surface. This process was introduced in [Chap. 4](#). It noted that a hard copy of an individual circuit was created in a reticle. This in turn is replicated on a full wafer mask (see [Fig. 4.14](#)). The mask-making process is detailed in [Chap. 9](#). A&E systems include use of a full mask or just the reticle to cover the wafer with the patterns in a step-by-step process (a stepper).

Nonmask systems have the required pattern coded into the exposure source. The pattern is directly “projected” to the photoresist layer without a mask as the exposing energy wave is “steered” to the right locations on the wafer surface.

The second subsystem is the exposure source. They include both optical and nonoptical sources ([Fig. 8.40](#)). Optical aligners use an ultraviolet light source, whereas nonoptical systems use exposure sources from other parts of the electromagnetic spectrum. The systems in use today were developed to keep pace with the reduction of feature size, increased circuit density, and required defect reductions of the ULSI era.

Alignment Selection Criteria
• Resolution Capability/ Limit • Alignment Accuracy • Contamination Level • Reliability • Productivity • Overall cost of ownership (COO)

FIGURE 8.40 Aligner selection criteria.

Aligners are selected and compared by several criteria ([Fig. 8.41](#)) that relate to their ability to produce the required images in a consistent and productive manner. Perhaps the most important parameter is the *resolution capability*, or the ability of the machine to produce a particular size image. The higher the resolution capability,

the better the machine. In addition to the resolution of the required image size, the aligner must be capable of placing the images in their correct position relative to each other. This performance parameter is called the *registration capability* of the aligner. These two factors must be performed over the entire wafer, a factor called *dimensional control*. The final performance factor is cost of ownership, which includes initial purchase cost, *wafer throughput* (the time required to load, align, expose, and unload the wafer), maintenance cost, and the up-time of the machine. These factors are discussed in [Chap. 15](#).

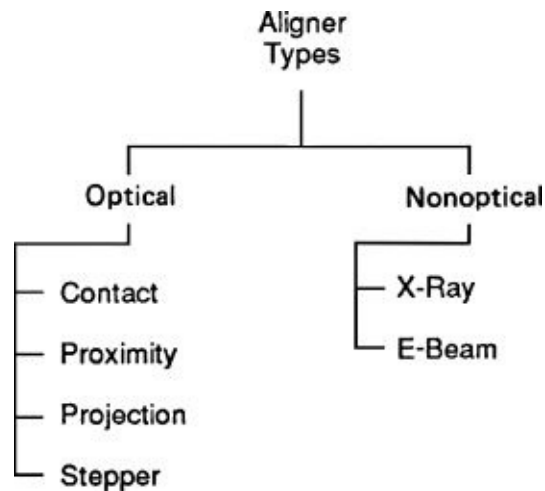


FIGURE 8.41 Table of aligner types.

So the options are a hard-pattern system (mask/reticle) or a direct write alignment. And there are options for exposure sources.

Exposure Sources

While aligners are very sophisticated machines, they operate on several basic optical principles. Consider producing a shadow image of a fork on a wall by shining a flashlight through a real fork held near the wall. By semiconductor standards, the image of the fork is quite imprecise. But there are ways to improve the image. One way would be to replace the flashlight with a narrower-wavelength light. The white light that comes from a flashlight contains many different wavelengths (colors), all mixed. At the edge of the fork, a phenomenon called *diffraction* takes place. Diffraction is the bending of the light rays at an opaque edge (or going through a narrow slit). The amount of bending depends on the wavelength, and, with many wavelengths (white light), there are many rays scattering from the edge and fuzzing up the image. Using a shorter-wavelength light, or a single-wavelength source, minimizes the diffraction. Another way to improve the image is to get all of the rays traveling along the same path. In normal white light, the rays are coming from the light bulb in many directions and leave the edge of the fork with many directions, again fuzzing the image. With mirrors and lenses, the light rays can be *collimated* into a band of parallel rays, thus improving image quality. The sharpness of the image and its dimensions are also influenced by the distance of the light behind the fork and the distance of the fork to the wall. Closing up both of these distances sharpens the image. These techniques (narrower or single wavelength exposure sources, collimated light, and strict control of the distances) are all used in aligners to produce the required images.

Exposure sources are chosen to create the required image size in conjunction with a specific photoresist (see “Photoresist Exposure Speed, Sensitivity, and Exposure Source”). The workhorse exposure source has been the high-pressure mercury lamp. It produces light in the ultraviolet (UV) range. The reduction of feature size has driven the development of improvements to the basic lamp and photo resists. To achieve greater definition, resists are tailored to respond to only narrow bands (lines) in the mercury lamp spectrum. This need also has resulted in the development of lamps and resists that operate in the shorter wavelengths of the spectrum. This region of the spectrum is known as the *deep ultraviolet* or DUV.

Other approaches to lower wavelength and higher-energy exposure sources are excimer lasers, X-rays, and electron beams. A more detailed discussion of the various exposure sources is provided in [Chap. 10](#).

Alignment Criteria

The first mask is aligned by positioning the y axis of the mask at a 90° angle to the major wafer flat on the wafer ([Fig. 8.42](#)). Subsequent masks are aligned to a previously patterned mask with the use of alignment marks (also called *targets*). These are special patterns ([Fig. 8.43](#)) located in an easily found position on the edge of each chip pattern or in the separation lines surrounding each chip on the wafer.

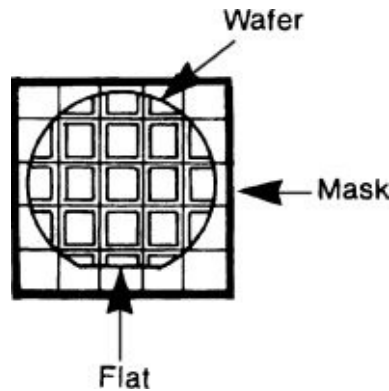


FIGURE 8.42 First mask or wafer alignment.

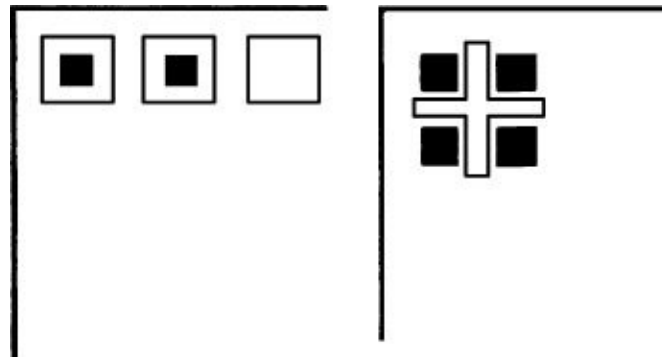


FIGURE 8.43 Types of alignment marks.

Alignment is accomplished by positioning a mark on the mask to a corresponding mark already on the wafer pattern from a previous photomasking step (see “Steppers”). Alignment marks become a permanent part of the chip surface through the etching process. They are then in place for the alignment of the next layer.

Alignment errors, called *misalignment*, fall into several categories ([Fig. 8.44](#)). A common one is a simple misplacement in the x-y directions. Another common misalignment is rotational, where one side of the wafer is aligned, but the patterns become increasingly misaligned across the wafer. A third rotational misalignment comes about when the die pattern is rotated on the mask or reticle.

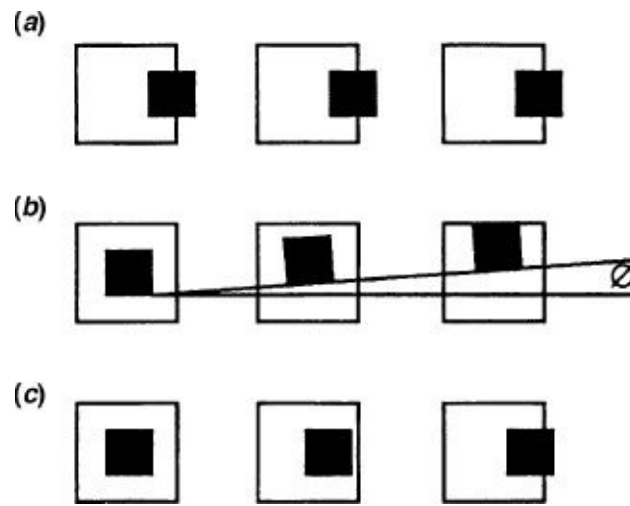


FIGURE 8.44 Misalignment types. (a) x direction, (b) rotational, and (c) runout.

Other misalignment problems associated with masks and stepper aligners are runout and run-in. These problems arise when the chip patterns are not formed on the mask on constant centers or are placed on the chip-off center. The result is that only a portion of the mask chip patterns can be properly aligned to the wafer patterns. The pattern becomes progressively misaligned across the wafer.

A rule of thumb is that circuits with micron or submicron feature sizes must meet registration tolerances of one-quarter to one-third the minimum feature. An *overlay budget* is calculated for the total circuit. It is the allowable accumulated alignment error for the entire mask set (see [Fig. 8.2](#)). For a 0.35- μm product, the allowable overlay budget is about 0.1 μm .⁵

Aligner Types

Contact Aligners

Until the mid-1970s, the contact aligner was the workhorse aligner of the semiconductor industry. The alignment part of the system uses a full wafer-size photomask positioned over a vacuum wafer chuck. The wafer is mounted on the chuck and viewed through a split-field objective microscope ([Fig. 8.45](#)). The microscope presents the operator with a simultaneous view of each side of the mask and wafer. The chuck is moved left, right, and/or rotated (x, y, and z movement) by manual controls until the wafer is aligned to the mask pattern.

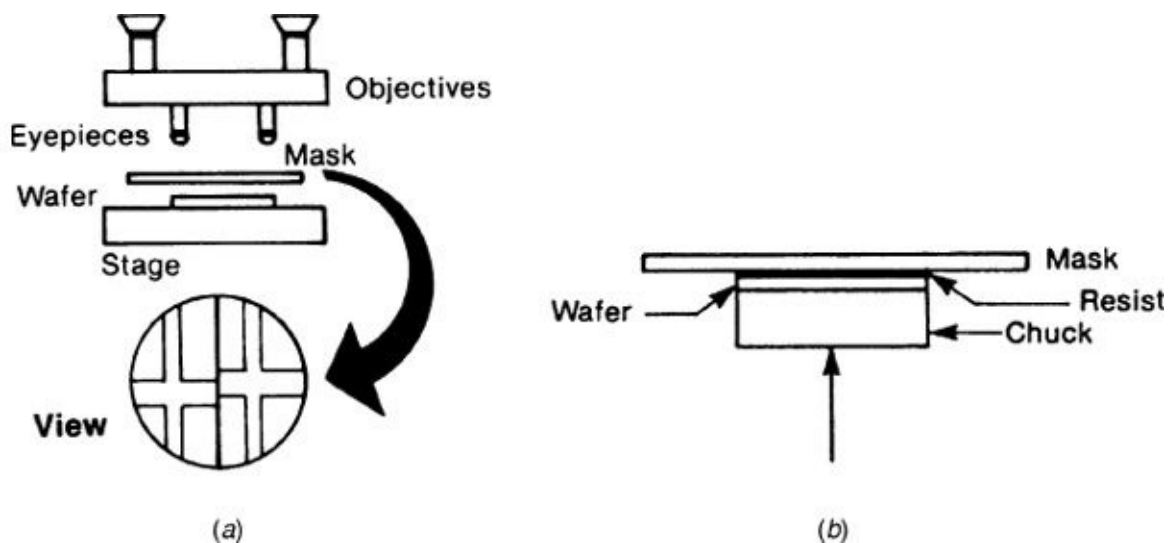


FIGURE 8.45 Contact aligner system. (a) Alignment stage and (b) contact stage.

Once the mask and wafer are aligned properly, the wafer chuck moves up on a piston, pushing the wafer into contact with the mask. Next, the collimated ultraviolet rays coming from a reflection and lens system pass through the mask and into the photoresist.

Contact aligners are used in production for discrete devices and circuits with SSI and MSI densities and feature sizes of approximately $5\text{ }\mu\text{m}$ and above. They also are used for flat-panel displays, infrared sensors, device packages, and multichip modules (MCMs).⁶ A contact aligner is capable of submicron imaging with the proper resist and a well-tuned process. Contact can damage the soft resist layer, the mask, or both. Dirt adhering to the clear portions of the mask blocks light during exposure. Epitaxial layer spikes on bipolar wafers can degrade the mask. Mask damage is so prevalent that masks have to be removed and discarded or cleaned every 15 to 25 exposures. Dirt between the mask and wafer will cause resolution problems in the immediate area of the piece of dirt. Alignment of larger diameter

wafers presents a light uniformity problem that causes image size variations and alignment problems.

Proximity Aligners

Proximity aligners were a natural evolution of contact aligners. The systems are essentially contact aligners but with mechanisms that hold the wafer either in near or soft contact with the mask. Sometimes proximity aligners are called *soft-contact machines*.

The performance of a proximity aligner is a tradeoff between resolution capability and defect density. With the wafer in soft contact with the mask, there is always some scattering of the light, which fuzzes the definition of the image in the resist. On the other hand, the soft contact also greatly reduces the number of defects associated with mask and resist damage. Even with the improved defect density, proximity aligners do not find much use in VLSI photomasking processes.

Scanning Projection Aligners

The end of contact aligners was foreseen for years, and development work was ongoing in the search for an alternative. The search centered on the concept of projecting ([Fig. 8.46](#)) the mask image onto the wafer surface, much as a slide (the mask) is projected onto a screen (the wafer). While simple in concept, the technique requires an excellent optical system to expose an entire wafer surface in one exposure. The problem was addressed with the introduction of the Perkin Elmer scanning projection aligner. This system avoided the problems of a full mask projection exposure in favor of a scanning technique that used a mirror system with a slit blocking part of the light coming from the light source. With this system, a new parameter, *scan speed*, became a parameter requiring control. They are called 1:1 aligners, since the image dimensions on the mask are the same size as the intended image dimensions on the wafer surface.

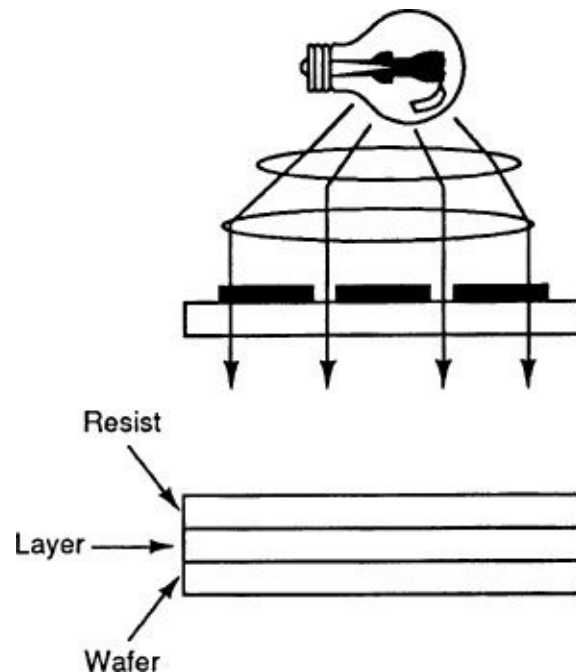


FIGURE 8.46 Concept of projection exposure.

Steppers

Scanning projection aligners were a great leap forward over contact aligners for production work, but they still had some limitations, such as alignment and overlay (registration) problems associated with full-size masks, image distortion, and mask-induced defects from dust and glass damage.

The next step was stepping images ([Fig. 8.47](#)) from a reticle directly onto the wafer surface—the same technique used to make masks. A reticle, carrying the pattern of one or several chips, is aligned and exposed, then is *stepped* to the next site and the process *repeated*. A reticle is of a higher quality than a full-size mask so fewer defects occur. There is better overlay and alignment, because each chip is individually aligned. The procedure of stepping allows precise matching of larger-diameter wafers. Other advantages are resolution improvements, because a smaller area is being exposed each time and a lessened vulnerability to dust and dirt. Some steppers are 1:1, that is, the image on the reticle is the same dimensions required on the wafer. Others use a reticle 5 to 10 times the final dimensions. These are called *reduction steppers*. Making an oversize reticle is easier, and any dirt and small glass distortions are reduced out of existence during the exposure ([Fig. 8.50](#)).

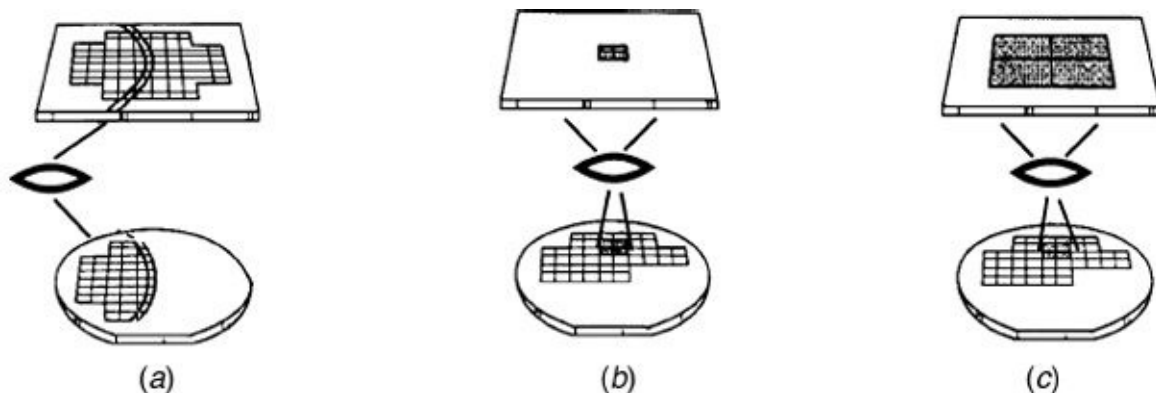


FIGURE 8.47 Projection imaging techniques. (a) Scan, (b) 1:1 step/repeat, and (c) reduction step/repeat.

A 5 $\times$ reticle is smaller and easier to make than a 10 $\times$ one. Also, a 5 $\times$ reticle projects a larger (up to 20 $\times$ 20 mm) field onto the wafer, resulting in faster wafer throughput ([Fig. 8.48](#)).² Field size is projected in the ITRS to grow to 9-in reticles with the capability of 25 $\times$ 50 mm field sizes.⁸

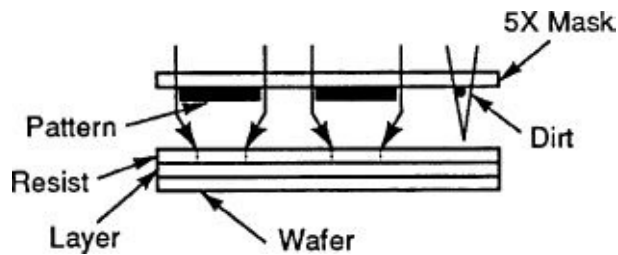


FIGURE 8.48 5× mask-pattern transfer.

The key to production use of steppers is automatic alignment systems. There is no way that an operator could individually align several hundred die on a wafer at a productive rate. Automatic alignment is accomplished by passing low-energy laser beams through alignment marks on the reticle and reflecting them off corresponding alignment marks on the wafer surface. The signal is analyzed, and information is fed to the x - y - z wafer chuck controls by a computer, which moves the wafer around until the wafer and reticle are aligned. The images are placed in the photoresist by sequentially exposing each die pattern across and down the wafer ([Fig. 8.49](#)).

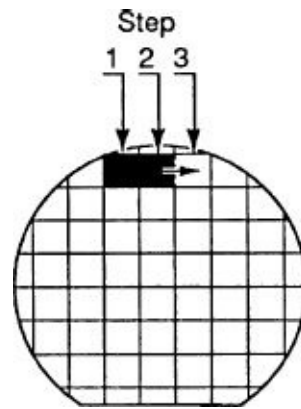


FIGURE 8.49 Step-and-repeat die alignment and exposure.

An alternative to laser signal control is a vision system. These systems use a camera to capture a vision of the die and compare it to a database. The wafer is moved until it and the mask or reticle image match the database.

Alignment systems using alignment marks ([Fig. 8.43](#)) are called off-axis alignment because the alignment mark is a reference but not part of the actual circuit pattern. Having off-axis alignment is another source of alignment errors. A more direct approach to alignment is through the lens (TTL). A TTL system looks directly at the wafer pattern. They use either the exposure light from a He-Ne or argon laser that bounces a signal off the wafer pattern and make the alignment adjustments automatically.⁹

Most production steppers have UV exposure sources with G-or I-line capabilities.

Steppers intended for small geometries are fitted with laser sources operating in the DUV range.⁹ The maintenance of the correct image size during the exposure part of the process requires tight humidity and temperature control. Most steppers are enclosed in an environmental chamber that controls these important parameters and keeps the wafers clean.

Step and Scan Aligners

Larger die sizes would normally require larger lens systems with larger fields of vision. Increasing the field of vision shortens the time of alignment and exposure. However, larger lenses become expensive. An alternative is a stepper with a smaller lens and the capability to scan the smaller field over the required area (see [Fig. 8.50](#)).

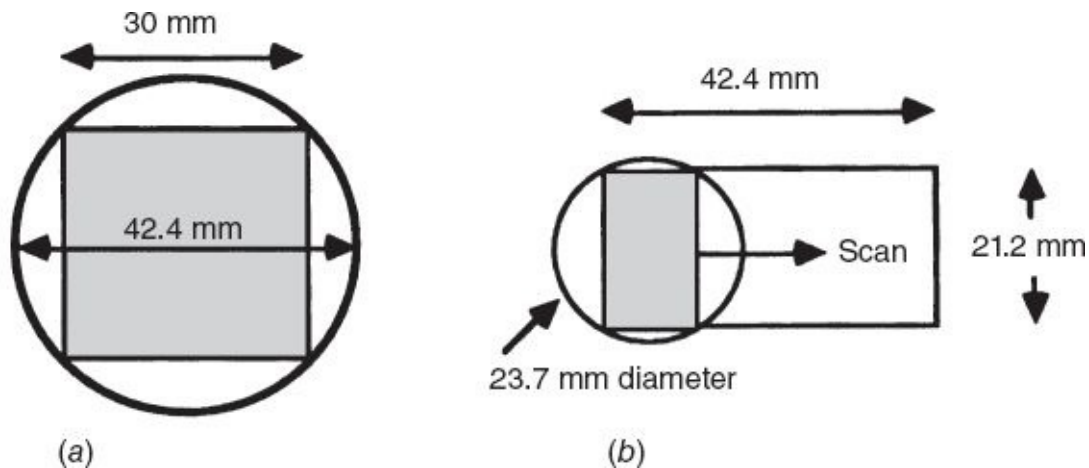


FIGURE 8.50 Step and scan comparison. (a) Step-and-repeat requires a 42.4-mm diameter lens field for 9-cm² die; (b) step-and-scan for same die requires a 23.7-mm lens field.

Other exposure sources and advanced alignment and exposure tools are explored in [Chap. 10](#).

Postexposure Bake

Standing waves are a problem that occurs with optical exposure and positive resists (see [Chap. 10](#)). One technique for minimizing the effect of standing waves is to bake the wafers after exposure. The bake method could be any one of the ones described earlier. The time-temperature specifications for the postexposure bake (PEB) are a function of the baking method, exposure conditions, and resist chemistry.

Electron Beam Aligners

Electron-beam lithography is a mature technology used in the production of high-quality masks and reticles ([Fig. 8.51](#)). The system consists of an electron source that produces a small-diameter spot and a “blanker” capable of turning the beam on and off. The exposure must take place in a vacuum to prevent air molecules from interfering with the electron beam. The beam passes through electrostatic plates capable of directing (or steering) the beam in the x - y direction on the mask, reticle, or wafer. This system is functionally similar to the beam-steering mechanisms of television sets. Precise direction of the beam requires that the beam travel in a vacuum chamber in which there is the electron beam source, support mechanisms, and the substrate being exposed.

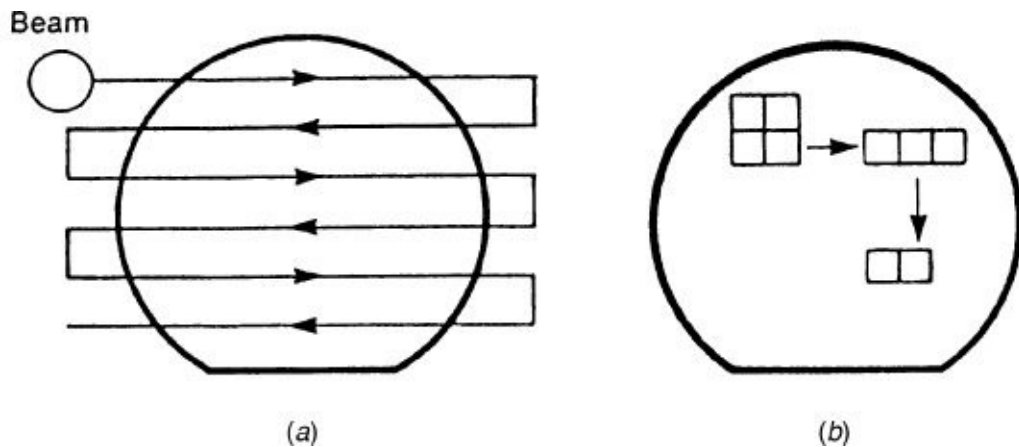


FIGURE 8.51 Electron beam scanning: (a) raster scan and (b) vector scanning.

Since the pattern required generates from the computer, there is no mask. The beam is directed to specific positions on the surface by the deflection subsystem and the beam turned on where the resist is to be exposed. Larger substrates are mounted on an x - y stage and are moved under the beam to achieve full surface exposure. This alignment and exposure technique is called *direct writing*.

The pattern is exposed in the resist by either raster or vector scanning ([Fig. 8.50](#)). Raster scanning is the movement of the electron beam side to side and down the wafer. The computer directs the movement and activates the blanker in the regions where the resist is to be exposed. One drawback to raster scanning is the time required for the beam to scan, since it must travel over the entire surface. In vector scanning, the beam is moved directly to the regions that have to be exposed (see [Fig. 8.51](#)). At each location, small square or rectangular shaped areas are exposed, building up the desired shape of the exposed area.

Mix and Match Aligners

Small-geometry imaging is expensive. Fortunately, only certain layers of a product mask set are critical enough to require the advanced imaging techniques. For advanced circuits, fully 50 percent of the layers may be noncritical.¹⁰ The other, less critical layers can be imaged with more established techniques such as projection scanners or less expensive steppers. Mix and match will probably be a feature of fab operations that use X-ray or e-beam technologies.

Advanced Lithography

The industry is careening along Moore's law toward a near future 5 nm node.¹¹ The basic processes described in [Chaps. 8](#) and [9](#) would not suffice to produce feature sizes much below the 200-nm node. Advancing beyond this milestone requires a whole host of improvements on the basic processes.¹² They include new resists, new exposure sources, improved masks, and more as addressed in [Chap. 10](#).

Review Topics

Upon completion of this chapter, you should be able to:

1. Sketch wafer cross-sections showing the basic ten-step photomasking process.
 2. Explain the reaction of positive and negative photo resists to light.
 3. Describe the correct resist and mask polarities required to produce holes and islands in wafer surface layers.
 4. Make a list of the major process options for each of the ten basic steps.
 5. Select from the list in objective 4 the processes used to pattern features in micron and submicron sizes.
 6. Describe the need for, and process steps used in, double masking, multilayer resist processing, and planarization techniques.
 7. Explain the use of antireflective coatings and contrast enhancement in the patterning of “small” feature sizes.
 8. List the optical and nonoptical methods used for alignment and exposure.
-

References

- [1.](#) Elliott, D. J., *Integrated Circuit Fabrication Technology*, McGraw-Hill, 1976, New York, NY:168.
- [2.](#) “Photoresists for Microlithography,” *Solid State Technology*, PennWell Publishing Company, Jun. 1993:42.
- [3.](#) Ibid., p. 71.
- [4.](#) Ibid., p. 116.
- [5.](#) Simon, K., “Abstract Alignment Accuracy Improvement by Consideration of Wafer Processing Impacts,” *SPIE Symposium on Microlithography*, 1994:35.
- [6.](#) Cromer, E., “Mask Aligners and Steppers for Precision Microlithography,” *Solid State Technology*, PennWell Publishing Company, Apr. 1993:24.
- [7.](#) Wolf, S., and Tauber, R. N., *Silicon Processing*, vol. 1, 2000, Lattice Press, Sunset Beach, CA:473.
- [8.](#) Singh, R., Vu, S., and Sousa, J., “Nine-Inch Reticles: An Analysis,” *Solid State Technology*, Oct. 1998:83.
- [9.](#) Fuller, G., *Optical Lithography, Handbook of Semiconductor Manufacturing Technology*, 2008, CRC Press, Hoboken, NJ:18–11.
- [10.](#) Cromer, E., “Mask Aligners and Steppers for Precision Microlithography,” *Solid State Technology*, PennWell Publishing Company, Apr. 1993:26.
- [11.](#) Shankland, S., “Moore’s Law: The Rule That Really Matters in Tech,” *CNET News*, Oct. 2012.
- [12.](#) Greeneich, J., “Mixing Critical and Noncritical Steppers,” *Solid State Technology*, PennWell Publishing Company, Oct. 1994:79.

CHAPTER 9

The Ten-Step Patterning Process— Developing to Final Inspection

Introduction

In this chapter, the basic patterning process steps used for photoresist developing through final inspection (steps 5 through 10 of the basic process) are explained. The end of the chapter examines the processes used for mask making and a discussion of alignment error budgets.

Development

After the wafer completes the alignment and exposure step, the device or circuit pattern is coded (*latent image*) in the photoresist as regions of exposed and unexposed resist ([Fig. 9.1](#)). The pattern is “developed” in the resist by the chemical dissolution of the unpolymersed resist regions. Development processes are designed to form a pattern with the exact dimensions designated during the circuit design process ([Fig. 9.1](#)). A problem resulting from a poor developing process is underdevelopment, which leaves the hole incompletely developed to the correct dimensions, or a coved sidewall. In some cases, the development will not be long enough (incomplete) and will leave a layer of resist in the hole. The third problem is overdevelopment, which removes too much resist from the image edge or top surface. High-aspect plug holes are a particular challenge with uniform hole diameter and cleaning being difficult due to the lack of fluid circulation in the deep hole.

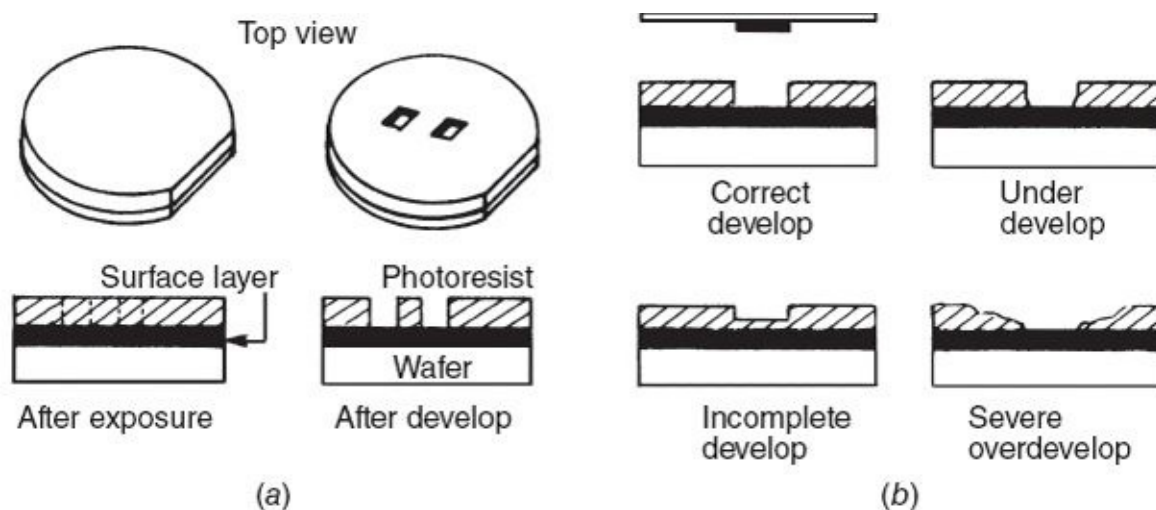


FIGURE 9.1 Photoresist development: (a) process; (b) problems.

Positive and negative resists have different developing characteristics and require different chemicals and processes ([Fig. 9.2](#)).

	Positive Resist	Negative Resist
Developer	Sodium hydroxide (NaOH) Tetramethyl ammonium hydroxide (TMAH)	Xylene Stoddard solvent
Rinse	Water (H ₂ O)	n-Butylacetate

FIGURE 9.2 Resist developer and rinse chemicals.

Positive Resist Development

After exposure, the intended pattern is coded in the positive photoresist as regions of polymerized resist (the starting condition) and unpolymerized resist (caused by the exposure). The two regions, polymerized and unpolymerized, have a dissolving rate difference of about 1:4. This means that, during the developing step, some resist is always lost from the polymerized region ([Fig. 9.3](#)). The use of developers that are too aggressive or that have overly long developing times may result in an unacceptable thinning of the resist film, which in turn may cause it to lift or break down during the etch step.

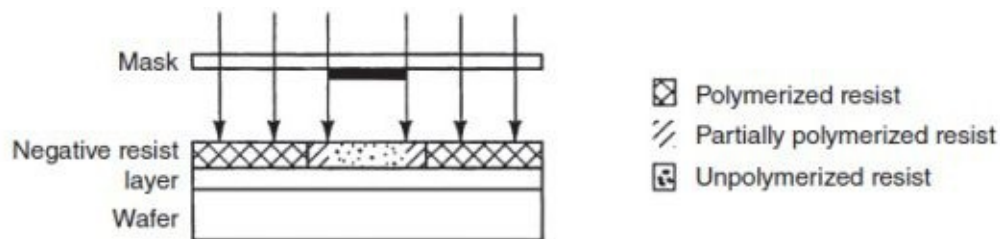


FIGURE 9.3 Transition region at resist image edge.

Two types of chemical developers are used with positive resist, alkaline-water solutions and nonionic solutions. The alkaline-water solutions can be sodium hydroxide or potassium hydroxide. Since both of these solutions contain mobile ionic contaminants, they are not desirable in processing sensitive circuits. Most positive-resist fabrication lines use a nonionic solution of tetramethylammonium hydroxide (TMAH). Sometimes a surfactant is added to break down the surface tension and make the solution more wettable to the wafer surface. The aqueous nature of positive developers makes them more environmentally attractive than the solvent developers required for negative resists.

Following the development step is a rinse to stop the development process and remove the development from the wafer surface. The rinse for positive resist is water that brings with it simpler processing, lower costs, and is environmentally beneficial.

The positive-resist developing process is more sensitive than negative processes.⁴ Factors influencing the outcome are the soft-bake time and temperature; degree of exposure; developer concentration; and time, temperature, and method of developing. The development process parameters are determined by matrix testing of all the variables. The effect on line width for a particular process is shown in [Fig. 9.4](#).

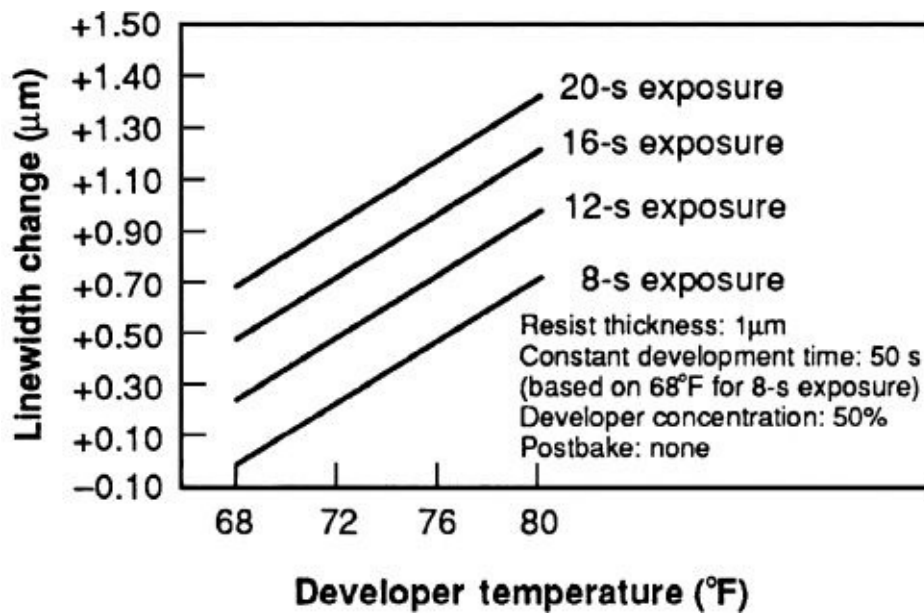


FIGURE 9.4 Developer temperature and exposure relationship versus line-width change.

Tight control of the development and rinse process is critical for dimensional control when using a positive resist. The rinse chemical for positive-resist developers is water. It serves the same role as negative-resist rinsers but is cheaper, safer to use, and allows easier disposal.

Negative Resist Development

The successful development of the image coded in the resist is dependent on the nature of the resist's exposure mechanisms. Negative resist, upon exposure to light, goes through a process of polymerization that renders the resist resistant to dissolution in the developer chemical. The dissolving rate between the two regions is high enough that little of the resist layer is lost from the polymerized regions. The chemical preferred for most negative-resist developing situations is xylene, which is also used as the solvent in negative-resist formulas.

The development step is done with a chemical developer followed by a rinse usually *n*-butyl acetate, because it neither swells nor contracts the resist. For wafers that have been patterned with a stepper, a milder-acting Stoddart solvent may be used.

Wet Development Processes

Several methods are used to develop resist films ([Fig. 9.5](#)). The selection of a method is dependent on the resist polarity, the feature size, defect density considerations, the thickness of the layer to be etched, and productivity.

Wet development methods

- Immersion
- Spray
- Puddle

FIGURE 9.5 Development methods.

Immersion

Immersion is the oldest development method. In its simplest form, the wafers, in a chemically resistant carrier, are immersed in a tank of the developer solution for a specific time before being transferred to a second tank of the rinse chemical ([Fig. 9.6](#)). The problems associated with such a simple wet procedure also define the development challenges:

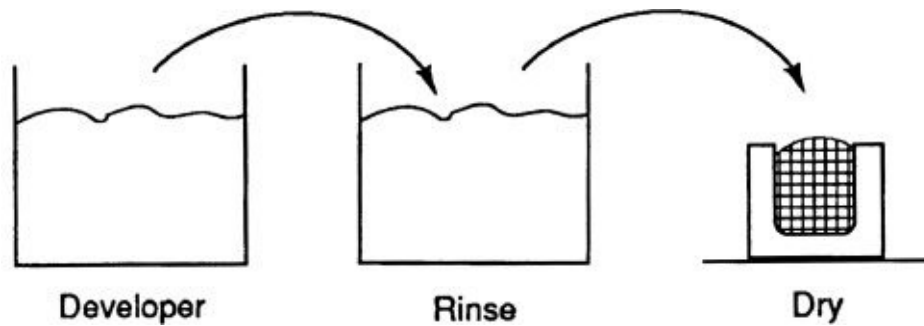


FIGURE 9.6 Immersion developer steps.

1. The surface tension of the liquids can prevent the chemicals from penetrating into small openings.
2. Partially dissolved pieces of resist can cling to the wafer surface.
3. The tanks can become contaminated as hundreds of wafers are processed through them.
4. The wafers can become contaminated as they are drawn through the liquid surface.
5. Developer chemicals (especially positive developers) can become diluted through use.
6. Frequent changing of solutions to eliminate problems 1, 2, and 3 raises the cost of the process.
7. Fluctuations in room temperature can cause changes in the developing rate of the solution.
8. The wafers have to be quickly transferred to a drying process step, which introduces a third step.

Additions are often made to the immersion tanks to improve the development process. Uniformity and penetration of small openings are aided by mechanical agitation of the bath by stirring or rocking mechanisms. A popular stirring system is a Teflon[®]-encapsulated magnet that is coupled to a rotating magnetic field outside the tank.

Agitation is also achieved by passing ultrasonic or megasonic waves through the liquid. The ultrasonic waves cause a phenomenon called *cavitation*. The energy in the waves causes the liquid to separate into tiny cavities that immediately collapse. The rapid generation and collapse of the millions of microscopic cavities create a uniformity of development and help the liquid penetrate into small openings. Sonic energy in the megasonic range (1 MHz) reduces a stagnant boundary layer that naturally clings to the wafer surface.² Uniform development rates are also enhanced by the addition of heaters and temperature controllers to the bath.

Spray Development

The preferred method of chemical development is by spray. In fact, spray processing is generally preferred over an immersion system for any wet process (clean, develop, etch) for a number of reasons. For instance, there is a major reduction in chemical use with spray systems. Process improvements include better image definition due to the mechanical action of the spray pressure in defining the resist edge and removing partially polymerized pieces of resist. Spray systems are always cleaner than immersion systems, because each wafer is developed (or etched or cleaned) only with fresh chemicals.

Spray processing is done in either single or batch systems. In the single-wafer configuration ([Fig. 9.7](#)), the wafer is clamped on a vacuum chuck and rotated while the developer and rinse are sequentially sprayed onto the surface. Drying takes place immediately after the rinse cycle by increasing the rotational speed of the wafer chuck. In appearance and design, a spray developer system for single wafers is the same as a resist spinner but plumbed for different chemicals. Single-wafer spray developers offer the advantage of track automation by integrating the developing and hard-bake processes. A major process advantage of these systems is the increased uniformity from the direct impingement of the chemical spray on the wafer.

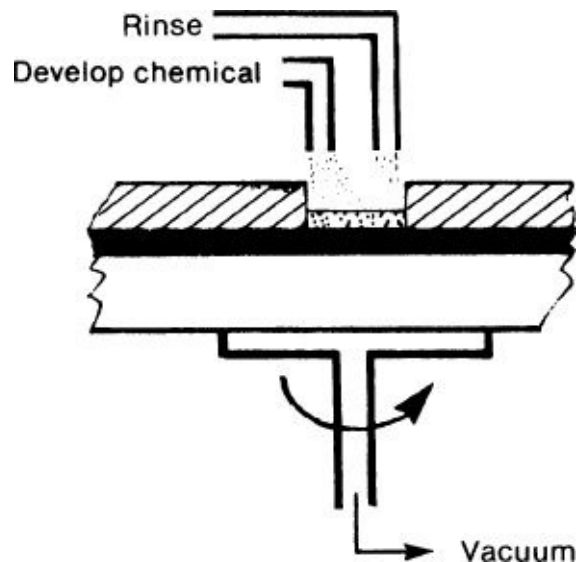


FIGURE 9.7 Spray development and rinse.

Positive resist development is more sensitive to temperature than negative resists. The problem is the phenomenon of rapid cooling of a fluid dispensed through an orifice under pressure. Called *adiabatic cooling*, it is the same phenomenon that

causes a can of pressurized household cleaner to cool during dispensing. Spray developers used for positive resist often have a heated wafer chuck or a heated spray nozzle to control the develop temperature. Other problems encountered with the spray development of positive resist are machine deterioration when alkaline developers are used and foaming as the water-based developer comes out of the nozzle under pressure. Batch developers come in two versions, single boat and multiple boat. These machines are the spin-rinse dryers described in [Chap. 7](#). As developers, they require additional plumbing to accommodate the developing chemicals but adapted for a development process. Batch developing systems, in general, are less uniform than direct-spray, single-wafer systems, because the spray does not impinge directly on each wafer surface, and temperature control for positive resist processing is more complicated.

Puddle Development

Spray development is very attractive for its uniformity and productivity. A process variation used to gain the advantages of spray for development of positive resist is the *puddle procedure*. This system uses a standard single-wafer spray unit. The difference between regular spray development and the puddle procedure is in the application of the developer chemical to the wafer. The process starts with the deposit of enough developer on the static wafer to cover the surface ([Fig. 9.8](#)). Surface tension holds the developer in a puddle on the wafer. The puddle sits there for some required time, usually on a chuck-heated wafer, causing the majority of the development to take place. Puddle development is, in effect, a single-wafer, front-side-only immersion process. After the required puddle time, the wafer surface is sprayed with more developer and rinsed, dried, and passed on to the next step.

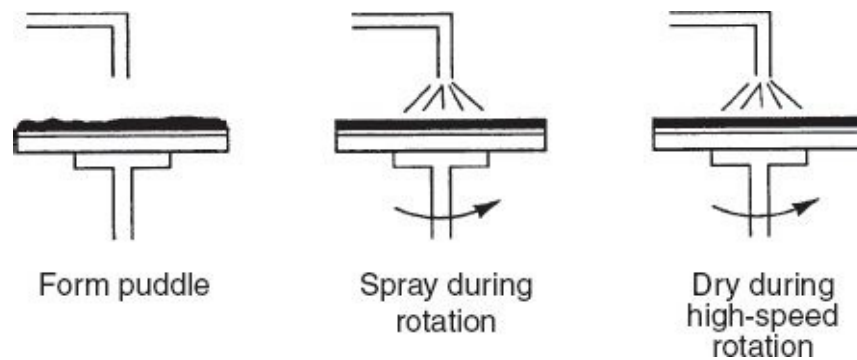


FIGURE 9.8 Puddle-spray development.

Plasma Descum

A particularly difficult form of incomplete development is a condition called *scumming*. The scum may be undissolved pieces of resist or dried developer³ left on the surface. The film is very thin and is difficult to detect with visual inspection. In reaction to this problem, advanced ultra large-scale integration (ULSI) lines with micron and submicron openings will remove (descum) the film from the wafers in an oxygen-rich plasma chamber after a chemical develop.

Dry (or Plasma) Development

The elimination of liquid processes has been a long-term industry goal. They are difficult to integrate into automated lines, and the chemicals are a substantial expense to purchase, store, control, and remove/dispose. One approach to replacing liquid chemical developers is to use a plasma-etch process. Dry plasma etch is a well-established process for etching wafer surface layers (see “Dry Etch”). In a plasma etcher, ions, energized by a plasma field, chemically dissolve (etch away) exposed layer surfaces. Dry resist development requires a photoresist chemistry that leaves either the exposed or unexposed portions of the resist layer readily removable by plasma-energized oxygen. In other words, one part of the pattern is oxidized off the wafer surface. A dry development process, called DESIRE, uses silylation and plasma O₂ and is described in [Chap. 10](#).

Hard Bake

Hard bake is the second heat treatment operation in the masking process. Its purpose is essentially the same as the soft-bake step: the evaporation of solvents to harden the resist. For hard bake, however, the goal is exclusively to achieve good adhesion of the resist to the wafer surface. As such, this step is sometimes called *pre-etch bake* or *prebake*.

Hard-Bake Methods

Hard bake is similar to soft bake in the equipment and methods used. Convection ovens, in-line and manual hot plates, infrared tunnel ovens, moving-belt conduction ovens, and vacuum ovens are all used for hard baking. The track systems are preferred for automated lines. See the “Soft Bake” section in [Chap. 8](#).

Hard-Bake Process

The exact time and temperature of the hard bake are determined much the same as in the soft-bake process. The starting point is the process recommended by the resist manufacturer. After that, the process is fine-tuned to achieve the adhesion and dimensional control required. Nominal hard-bake temperatures are from 130°–200°C for 30 min in a convection oven. Temperatures and times vary in the other methods. The minimum temperature is set to achieve good adhesion of the resist-image edge to the surface. The heat-caused adhesion mechanism is dehydration and polymerization. The heat drives water out of the resist, at the same time further polymerizing it, thereby increasing its etch-resistant properties.

The upper temperature limit of the hard bake is set by the flow point of the resist. Resist is a plastic-like material that softens and flows when heated ([Fig. 9.9](#)). When the resist flows, the image dimensions are changed. Extreme flow exhibits itself as a series of fringe lines around the image. The fringes are an optical effect from the slope left in the resist after the flow.

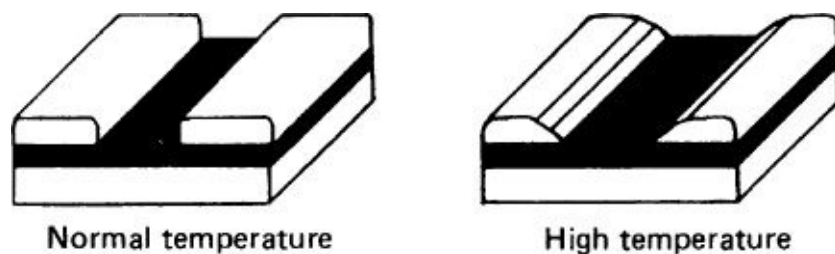


FIGURE 9.9 Resist flow at high temperature.

Hard bake takes place either immediately after the developing step or just before the etching step, as shown in [Fig. 9.10](#). In most production situations, the hard bake is performed in a tunnel oven that is in-line with the developer. When this procedure is used, it is important that the wafers be stored in a nitrogen atmosphere and/or be processed through the develop inspection step as quickly as possible to prevent the reabsorption of water into the resist film.

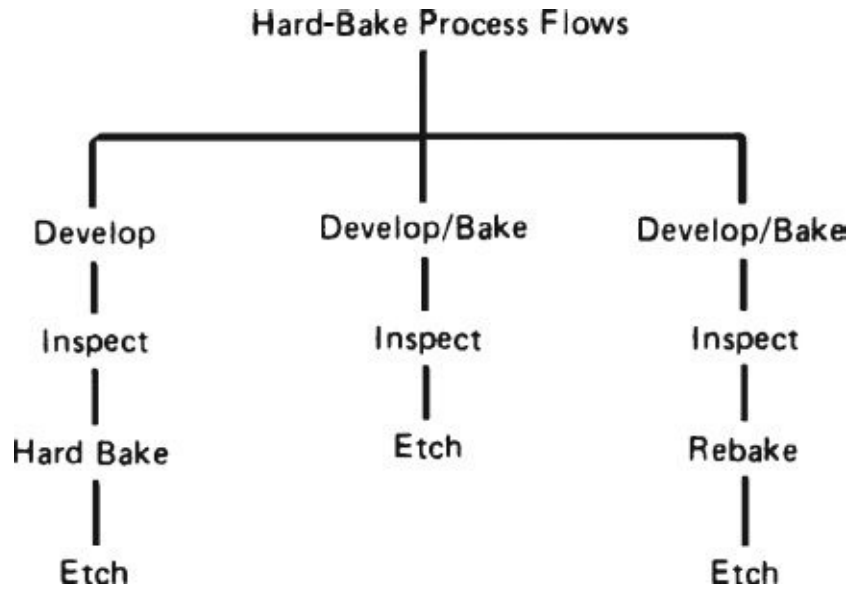


FIGURE 9.10 Hard-bake process flow options.

A goal of process engineering is to have as many common processes as possible. For hard baking, that is sometimes difficult due to the different adhesion characteristics of various wafer surfaces. The more difficult surfaces, such as aluminum-and phosphorusdoped oxides, are sometimes given a higher-temperature hard bake or a second hard bake in a convection oven just prior to being etched.

Develop Inspect

The first quality check in the photomasking process is performed after the development and baking steps. It is appropriately called *develop inspect* or simply DI. The purpose of the inspection is to identify wafers that have a low probability of passing the final masking inspection, provide process performance and process control data, and identify wafers for rework.

This inspection yield (that is, the number of wafers that pass this first quality check) is not factored into the overall yield formula, but it is a much-watched yield for two principal reasons. The critical nature of the photomasking process to the functioning of the circuit has been emphasized throughout this text. Develop inspect is the first chance to measure the performance of the photomasking process. The second importance of develop inspect is related to the two types of rejects made at this inspection. First, some of the wafers will have problems from previous steps that prevent their continuation in the process. These wafers are rejected at develop inspect and discarded. Other wafers have problems associated specifically with the quality of the pattern in the resist film. These wafers can sometimes be reworked ([Fig. 9.11](#)) by removal of the photoresist and reinsertion into the patterning process. This is one of the few places in the entire fabrication process where a general rework of mistakes is possible, since no permanent changes have been made to the wafer.

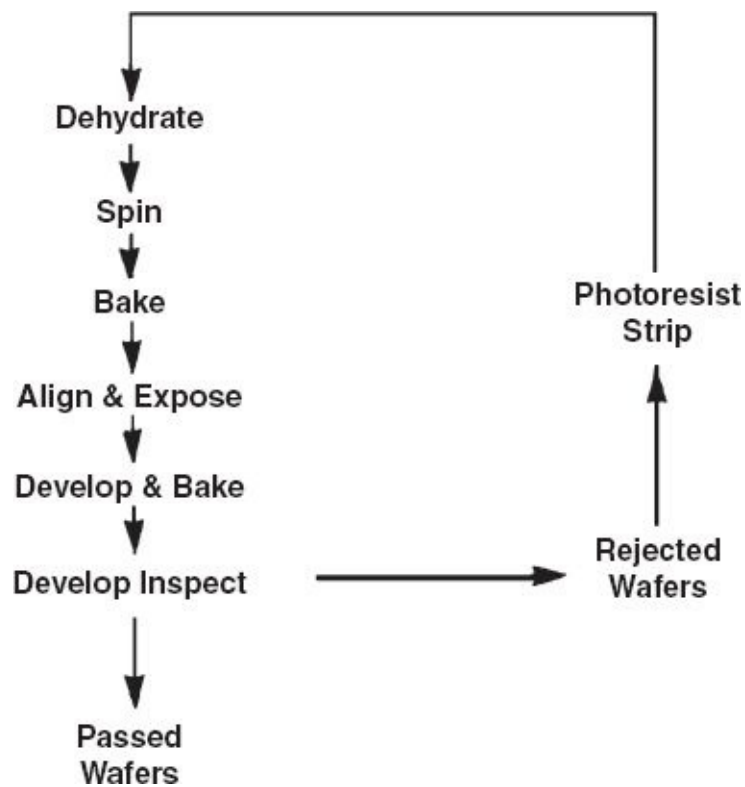


FIGURE 9.11 Rework process loop.

Wafers sent back into the masking process are called *reworks* or *redos*. The goal is to keep the rework rate as low as possible—certainly under 10 percent and preferably less than 5 percent. Experience has shown that wafers that have gone through a masking rework have a lower wafer-sort yield at the end of the fabrication process. Reworking causes adhesion problems, and the additional handling can result in contamination and breakage.

The develop inspect yield and rework rate vary from mask level to mask level. In general, the first levels in the masking sequence have wider feature sizes, flatter surfaces, and lower density, all of which make for a higher yield out of the mask step. By the time the wafers are at the critical contact and metal masks, the rework rate tends to rise.

Develop Inspect Reject Categories

In general, there are six primary categories of wafer problems that occur at both develop inspect and final inspect. There are:

- Out of spec pattern dimensions (critical dimension measurements)
- Misaligned patterns
- Surface contamination
- Holes or scratches in the photoresist
- Stains or other surface irregularities
- Patterns with distorted shapes

Develop Inspect Methods

Descriptions of the inspection equipment are presented in [Chap. 14](#).

Automatic Inspection

As die size got larger, dimensions got smaller, and processes became more numerous and more sophisticated, the older, and relatively slow manual inspection (see below), became nonproductive. Automatic inspection systems designed to detect surface and pattern distortion problems are the inspection system of choice for both off-line and on-line inspections. The systems are described in [Chap. 14](#). Automatic systems offer the prospect of higher amounts of data, and this in turn enables the process engineers to characterize and control the processes. They also are consistent, whereas humans vary in abilities and fatigue doing repetitive tasks. However, some process defects such as very small particles and problems from previous patterning steps can escape detection with automated systems. Optical microscopes are still useful for analysis.

Besides the increased productivity with automatic inspection systems, there are the issues of accuracy. As mentioned pattern line widths are getting smaller and in denser patterns smaller defects become killer defects. Scanning electron microscopes (SEMs) can measure smaller dimensions and detect smaller particles and surface defects. Atomic force microscopes measure surface flatness and detect irregularities. X-ray spectroscopy is used for contamination detection. The goal is in-line inspections of all wafers but some techniques require off-line analysis as well.

Some parameters, such as critical dimensions are measured on test patterns designed into the chip design. And the complexity of inspecting for defects and other surface problems is complicated as more layers are added to the wafer. Often blank “monitor” wafers will be included in the processing batch. Defects and contamination introduced during the specific process steps are more easily detected and measured on monitor wafers.

Manual Inspection

The flow diagram in [Fig. 9.12](#) shows a typical manual develop inspection sequence. The first step is a naked-eye inspection of the wafer surface. The wafer is viewed under a collimated white light or a high-intensity ultraviolet light. The wafer is viewed at an angle in the light beam. This method is surprisingly effective in showing film thickness irregularities, gross developing problems, scratches, and contamination, especially stains.

Step	Looking For:	Methods
1.	Stains/ gross contamination	Naked eye/UV light
2.	Stains/contamination Pattern irregularities Missalignment	100-400X microscope/ SEM/ AFM/ Automatic Inspection Systems
3.	Critical dimensions	microscope/SEM/ AFM

FIGURE 9.12 Order of develop inspection.

As the die density goes up, the individual parts also get smaller, which in turn requires a higher magnification to see them. The increasing magnification narrows the field of view, which in turn increases the time for an operator to inspect a wafer. The time required to statistically sample a large-diameter, low-defect wafer is prohibitive. Often, the microscopes will have motorized or programmable stages that automatically go to the inspection areas on the wafer.

Causes for Rejecting at the Develop Inspection Stage

There are many reasons why a wafer can be rejected or sent for rework at the develop inspection step. Generally, the only defects looked for are those added to the wafer during the current photomasking step. Defects from previous masking steps are generally overlooked under the theory that every wafer is passed on with some defects or problems and that the wafer arrived at the current step with acceptable quality. If there is a serious problem with a wafer that somehow escaped previous inspections, it is pulled from the batch.

The inspection is generally on a *first-fail basis*. The inspection continues until a rejectable level is reached on the wafer and the wafer is identified as a reject. The information for each wafer is logged for data accumulation and analysis. Automatic and semiautomatic optical inspection stations have electronic memories for accumulating and correlating this reject data.

Typical reject causes at develop-inspect stage are:

- Broken wafer
- Scratch
- Contamination
- Pinholes in resist
- Misaligned pattern
- Bridging
- Lifting resist
- Underexposure
- Incomplete development
- No resist
- Resist flowing
- Wrong mask
- Critical dimension(s)

Most of the causes for rejects have been discussed. One problem not discussed is bridging ([Fig. 9.13](#)). It is a condition in which two patterns are connected (bridged) by a thin layer of photoresist, usually at the metal layer. If passed on to the etch step, the photoresist bridge results in an electrical short between the patterns. Bridging comes from an overexposure, poor mask definition, or a resist film that is too thick. Bridging is a particularly vexing problem as patterns get closer together.

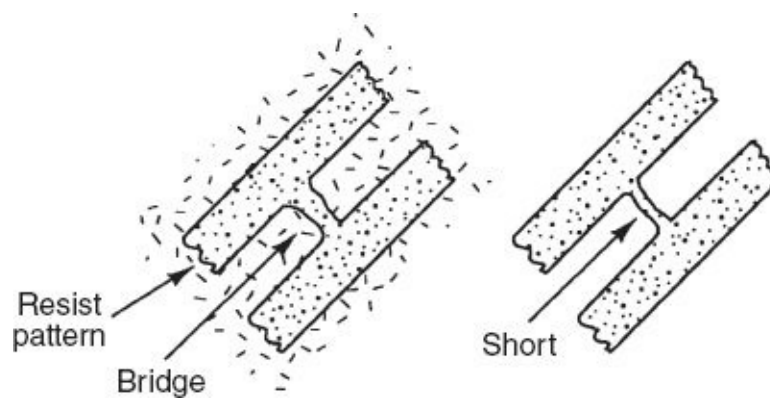


FIGURE 9.13 Bridged conduction lines.

Etch

At the completion of the develop inspect step, the mask (or reticle) pattern is defined in the photoresist layer and is ready for etch. In the etch step the image is permanently transferred into the surface layer on the wafer. Etching is the process of removing the top layer(s) from the wafer surface through the openings in the resist pattern.

Etching processes fall into two main categories: wet and dry (see [Fig. 9.20](#)). The primary goal of each is an exact transfer of the image from the mask or reticle into the wafer surface. Other etch process goals include uniformity, edge profile control, selectivity, cleanliness, and minimizing cost of ownership (COO) of the equipment and process stations.

Wet Etching

The historic method of etching has been by immersion techniques using wet etchants. The procedure is similar to the preoxidation clean-rinse-dry process ([Chap. 7](#)) and immersion development. Wafers are immersed in a tank of an etchant for a specific time, transferred to a rinse station for acid removal, and transferred to a station for final rinse and a spin-dry step. Wet etching is used for products with feature sizes greater than 3 μm . Below that level, the control and precision needed requires dry-etching techniques.

Etching uniformity and process control are enhanced by the addition of heaters and agitation devices, such as stirrers or ultrasonic and megasonic waves, to the immersion tanks.

Wet etchants are selected for their ability to uniformly remove the top wafer layer without attacking the underlying material (good selectivity).

Etch-time variability is a process parameter influenced by temperature variations as the boat and wafers come to temperature equilibrium in the tank and the continued etching action as the wafers are transferred to a rinse tank. Generally, the process is set at the shortest time compatible with uniform etching and high productivity. The maximum time is limited to the amount of time the resist will continue to adhere to the wafer surface.

Etch Goals and Issues

The exactness of the image transfer in etching is dependent on several process factors. They include incomplete etch, overetching, undercutting, selectivity, and anisotropic or isotropic etching of the sidewalls.

Incomplete Etch

Incomplete etch is a situation in which a portion of the surface layer remains in the pattern hole or on the surface ([Fig. 9.14](#)). The causes of incomplete etch are a shortened etch time, the presence of a surface layer that slows the etching, or an uneven surface layer that results in incomplete etch in the thicker portions of the wafer. If wet-chemical etching is used, a lowered temperature or weak etch solution will cause incomplete etch. If dry-plasma etching is used, a wrong gas mixture or an improperly operated system can cause the same effect.

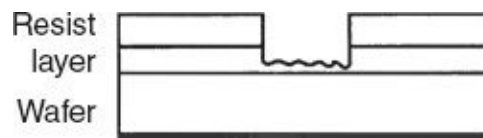


FIGURE 9.14 Incomplete etch.

Overetch and Undercutting

The opposite condition to incomplete etch is overetch. In any etch process, there is always some degree of overetch planned. This allows for thickness variations in the surface layer. Planned overetch also allows for the etch to break through any slow-etching layers on the top surface.

The ideal etch leaves vertical sidewalls in the layer ([Fig. 9.15](#)). Etch techniques that produce this ideal result are said to be *anisotropic*. However, the etchant actually removes material in all directions. This phenomenon is called *isotropic* etching. Etching takes place at the top of the layer for the entire time it takes to reach the bottom of the layer. The result is a sloped side. This action, is also called *undercutting* ([Fig. 9.16](#)), because the surface layer is undercut below the resist edge. An ongoing goal of the etch step is the control of undercutting to an acceptable level. Circuit layout designers take undercutting into account when planning the circuit. Adjacent patterns must be separated a certain distance to prevent shorting. The amount of undercutting must be calculated when the pattern is designed. The anisotropic etching available with plasma etching is preferred for advanced circuits. Reducing undercutting allows denser circuits.

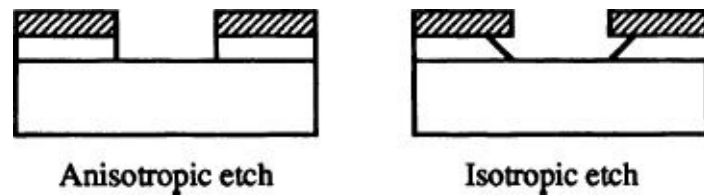


FIGURE 9.15 Anisotropic and isotropic etch.

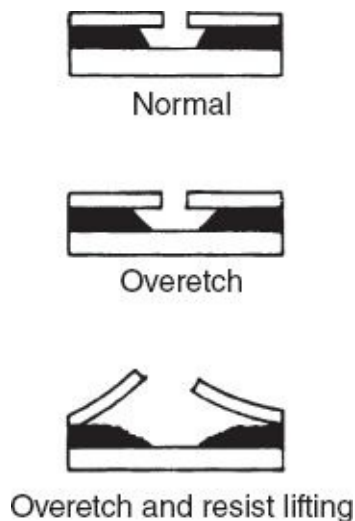


FIGURE 9.16 Degrees of undercutting.

Severe undercutting (or overetch) takes place when the etch time is excessive, the etching temperature is too high, or the etch mixture is too strong. Undercutting is also present when the adhesion bond between the photoresist and the wafer surface is weak. This is a constant worry, and the purpose of the dehydration, prime, soft-bake, and hard-bake steps is to prevent this type of failure. Failure of the resist bond at the edge of the etch hole can result in severe undercutting. If the bond is very poor, the resist can lift from the wafer surface, causing catastrophic undercutting.

Selectivity

Another goal in the etch step is the preservation of the surface underlying the etched layer. If the underlying surface of the wafer is partially etched away, the physical dimensions and electrical performance of the devices are changed. The property of the etch process that relates to preservation of the surface is *selectivity*. It is expressed as the ratio of the etching rate of the layer material to the etch rate of the underlying surface. Oxide or silicon selectivity varies from 20 to 40,⁴ depending on the etching method. High selectivity implies little or no attack of the underlying surface. Good selectivity becomes a problem in etching small contacts with aspect ratios greater than 3:1.⁵ Selectivity also applies to the removal of the photoresist. This is more of a consideration in dry-etch processes. While the top layer(s) is being removed by the etchant, some photoresist is also being removed. The selectivity factor must be high enough to ensure that the surface does not lose its protective photoresist layer before the etch is complete.

Wet-Spray Etching

Wet-spray etching offers several advantages over immersion etching. Primary is the added definition gained from the mechanical pressure of the spray.² Spray etching also minimizes contamination from the etchants. From a process control point of view, spray etching is more controllable, since the etchant can be instantly removed from the surface by switching the system to a water rinse. Single-wafer spinning-chuck spray systems offer considerable process uniformity advantages.

Disadvantages to spray etching are system cost, safety considerations associated with caustic etchants in a pressurized system, and the requirement of etch-resistant materials to prevent the deterioration of the machine. On the plus side, spray systems are usually enclosed, which adds to worker safety.

Batch-immersion etching, while featuring high productivity, also does not meet uniformity requirements for small feature sizes and/or larger-diameter wafers. Single wafer module spray tools with robot loading and unloading systems overcome the limitations of batch-immersion etching (see [Chap. 7](#)). They provide the needed control of the chemical composition, timing and uniformity of the etch. The common chemicals used to etch different layers in silicon technology are discussed in the following sections. [Figure 9.19](#) is a table of common semiconductor films and their common etchants.

Silicon Wet Etching

Silicon layers are typically etched with a solution of nitric and hydrofluoric (HF) acids mixed in water. The formula becomes an important factor in control of the etch. In some ratios, the etch has an exothermic reaction with the silicon. Exothermic reactions are those that produce heat, which in turn speeds up the etch reaction, which in turn creates more heat, and so on, resulting in an uncontrollable process. Sometimes, acetic acid is mixed with the other ingredients to control the exothermic reaction.

Some devices require the etching of a trough or trench into the silicon surface. The etch formula is adjusted to make the etch rate dependent on the orientation of the wafer. $\langle 111 \rangle$ -oriented wafers etch at a 45° angle, whereas $\langle 100 \rangle$ -oriented wafers etch with a “flat” bottom.⁶ Other orientations result in different-shaped trenches. Polysilicon films are also etched with the same basic formula.

Silicon Dioxide Wet Etching

The most common etched layer is a thermally grown silicon dioxide. The basic etchant is hydrofluoric acid, which has the advantage of dissolving silicon dioxide without attacking silicon. However, full-strength HF has an etch rate of about 300 Å/s at room temperature.⁶ This rate is too fast for a controllable process (a 3000-Å layer would etch in only 10 s).

In practice, the HF (assay of 49 percent) is mixed with water or ammonium fluoride and water. The ammonium fluoride (NH_4F) acts as a buffer to the unwanted generation of hydrogen ions which accelerate the etch rate. These solutions are known as *buffered oxide etches* or BOEs. They are mixed in different strengths to create reasonable etch times for the particular oxide thickness (Fig. 9.17). Some BOE formulas include a wetting agent (a surfactant such as Triton X-100 or equivalent) to reduce the surface tension of the etch, allowing it to uniformly penetrate into smaller openings.

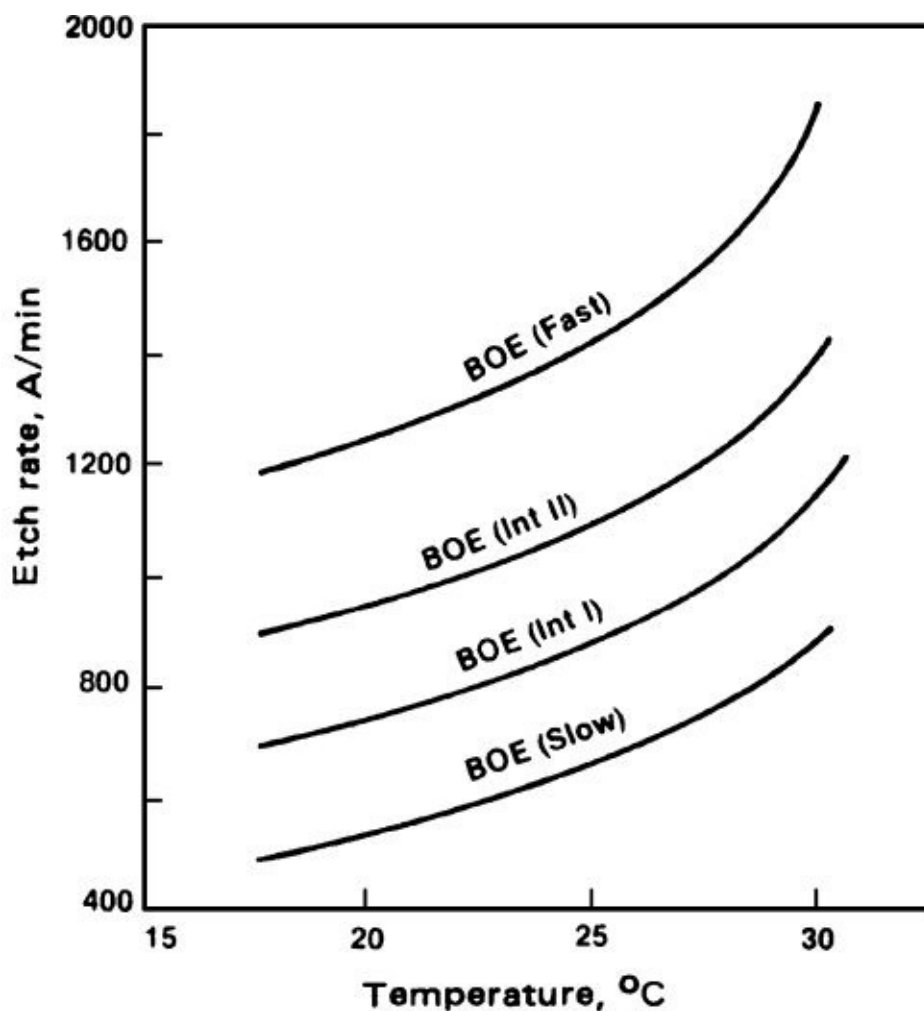


FIGURE 9.17 Etch rate versus temperature for BOEs.

Overetching that exposes the silicon wafer surface can cause surface roughing. The silicon surface can become roughed through etching when exposed to the OH^- ion during the HF process.

Aluminum-Film Wet Etching

Selective etching solutions for aluminum and aluminum alloy layers are based on phosphoric acid. An unfortunate by-product of the reaction of aluminum and phosphoric acid is tiny bubbles of hydrogen, as shown in the reaction in [Fig. 9.18](#). These bubbles cling to the wafer surface and block the etch action. The result is either bridges of aluminum that can cause electrical shorts between adjacent leads or spots of unwanted aluminum, called *snowballs*, left on the surface.

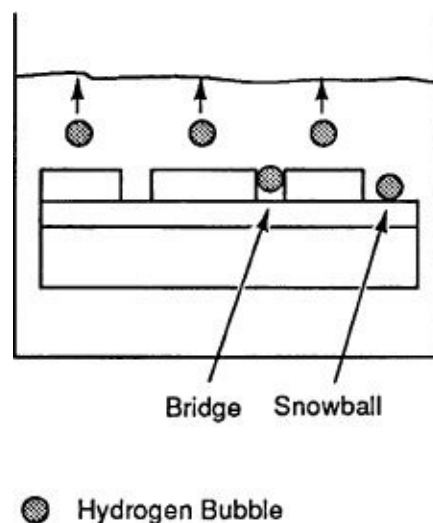


FIGURE 9.18 Hydrogen bubble blockage etchant.

Neutralization of this problem is accomplished by use of an aluminum etching solution that contains phosphoric acid, nitric acid, acetic acid, water, and wetting agents. A typical solution of the active ingredients (less wetting agent) is 16:1:1:2.

In addition to the special formulas, a typical aluminum etch process will include wafer agitation by stirring or moving the wafer boat up and down in the solution. Sometimes, ultrasonic or megasonic waves are used to collapse and move the bubbles around.

Deposited-Oxide Wet Etching

One of the final layers on a wafer is a silicon dioxide passivation film deposited

over the aluminum metallization pattern. These films are known as *vapox* or *silox* films. While the chemical composition of the films is that of silicon dioxide, the same as thermally grown silicon dioxide, they require a different etch solution. The difference is in the selectivity required of the etchant.

The usual etchant for silicon dioxide is a BOE solution. Unfortunately, the BOE attacks the underlying aluminum pads, causing bonding problems in the packaging process. This condition is called *brown*, or *stained, pads*. The preferred etchant for this layer is a solution of ammonium fluoride and acetic acid mixed in a ratio of 1:2.

Silicon Nitride Wet Etching

Another compound favored for the passivation layer is silicon nitride. It is possible to etch this layer by wet chemical means, but it is not as easy as for the other layers. The chemical used is hot (180°C) phosphoric acid. Because the acid evaporates rapidly at this temperature, the etch must be done in a closed reflux container equipped with a cooled lid to condense the vapors ([Fig. 9.19](#)). The major problem is that photoresist layers do not stand up to the etchant temperature and aggressive etch rate. Consequently, a layer of silicon dioxide or some other material is required to block the etchant. These two factors have led to the use of dry-etching techniques for silicon nitride.

	Common Etchant	Etch Temp	Rate A/min	Method
SiO ₂	HF & NH ₄ F (1:8)	Room	700	Dip & wetting agent predip
SiO ₂	HF & NH ₄ F (1:8)	Room	700	Dip & wetting agent predip
SiO ₂ (Vapox)	Acetic Acid & NH ₄ F(2:1)	Room	1000	Dip
Aluminum	H ₃ PO ₄ : 16 HNO ₃ : 1 Acetic : 1 H ₂ O : 2 Wetting agent	40–50°C	2000	a) Dip & agitation b) Spray
Si ₃ N ₄	H ₃ PO ₄	150–180°C	80	Dip
POLYSi	HNO ₃ : 50 H ₂ O : 20 HF : 3	Room	1000	Dip

FIGURE 9.19 Summary of wet-etching process.

Vapor Etching

Vapor etching is the exposure of the wafer to etchant vapors. HF is the most common. Advantages are continued replenishment of the etchant at the surface with fresh vapors and instant etch termination. Keeping the toxic vapors contained in the system is a safety concern.

Dry Etch

The limits of wet etching for small dimensions have been mentioned in the previous section. For review, they include the following:

1. Wet etching is limited to pattern sizes of 2 μm or larger.
2. Wet etching is isotropic, resulting in sloped sidewalls.
3. A wet-etch process requires rinse and dry steps.
4. The wet chemicals are hazardous and/or toxic.
5. Wet processes represent a contamination potential.
6. Failure of the resist-wafer bond causes undercutting.

These considerations have led to the use of dry-etch processes for the definition of small feature sizes on advanced circuits. [Figure 9.20](#) provides an overview of the dry etching techniques used.

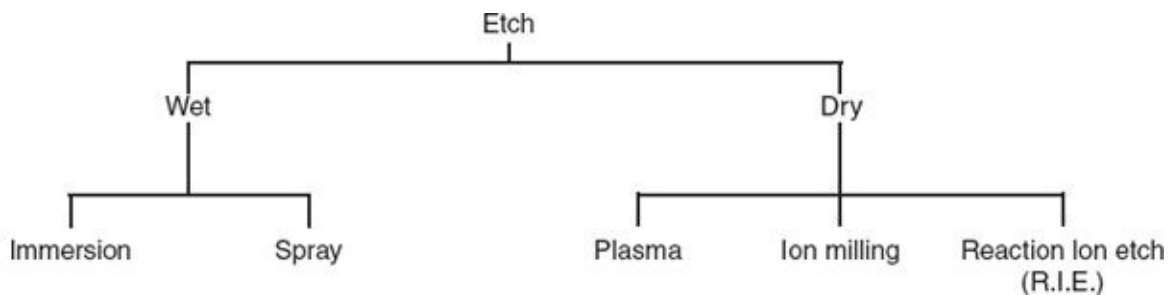


FIGURE 9.20 Guide to etch methods.

Dry etching is a generic term that refers to the etching techniques in which gases are the primary etch medium, and the wafers are etched without wet chemicals or rinsing. The wafers enter and exit the system in a dry state. There are three dry-etching techniques: plasma, ion beam milling (etching), and reactive ion etch (RIE).

Plasma Etching

Plasma etching, like wet etching, is a chemical process but uses gases and plasma energy to cause the chemical reaction. Comparison of silicon dioxide etching in the two systems illustrates the differences. In wet etching of silicon dioxide, the fluorine in the BOE etchant is the ingredient that dissolves the silicon dioxide, converting it to water-rinsable components. The energy required to drive the reaction comes from the internal energy in the BOE solution or from an external heater.

A plasma etcher requires the same elements as a wet etch: a chemical etchant and an energy source. Physically, a plasma etcher consists of a chamber, vacuum system, gas supply, an end-point detector, and a power supply ([Fig. 9.21](#)). The wafers are loaded into the chamber, and the pressure inside is reduced by the vacuum system. After the vacuum is established, the chamber is filled with the reactive gas. For the etching of silicon dioxide, the gas is usually CF_4 mixed with oxygen. A power supply creates a radio frequency (RF) field through electrodes in the chamber. The field energizes the gas mixture to a plasma state. In the energized state, the fluorine attacks the silicon dioxide, converting it into volatile components that are removed from the system by the vacuum system.

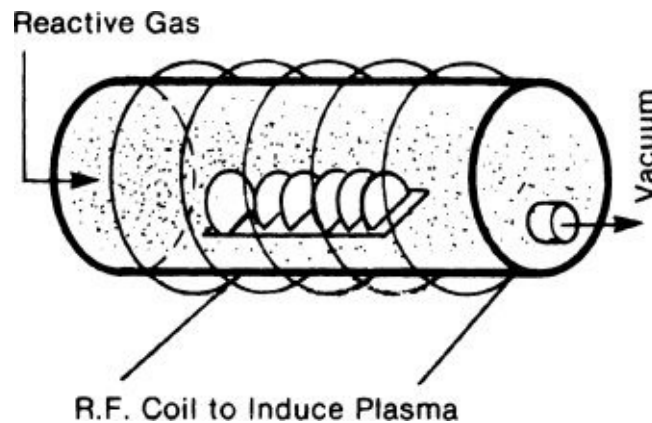


FIGURE 9.21 Barrel plasma etch.

Planar Plasma Etching

For more precise etching, plasma planar systems are used. These systems contain the basic components of the barrel system, but the wafers are placed on a grounded pallet under the RF electrode (Fig. 9.22). Etching takes place with the wafers actually in the plasma field. The etching ions are more directional than those in a barrel system, resulting in a more anisotropic etch. Almost vertical sidewalls are possible with plasma etch. Etching uniformity is increased with the rotation of the wafer pallet in the system.

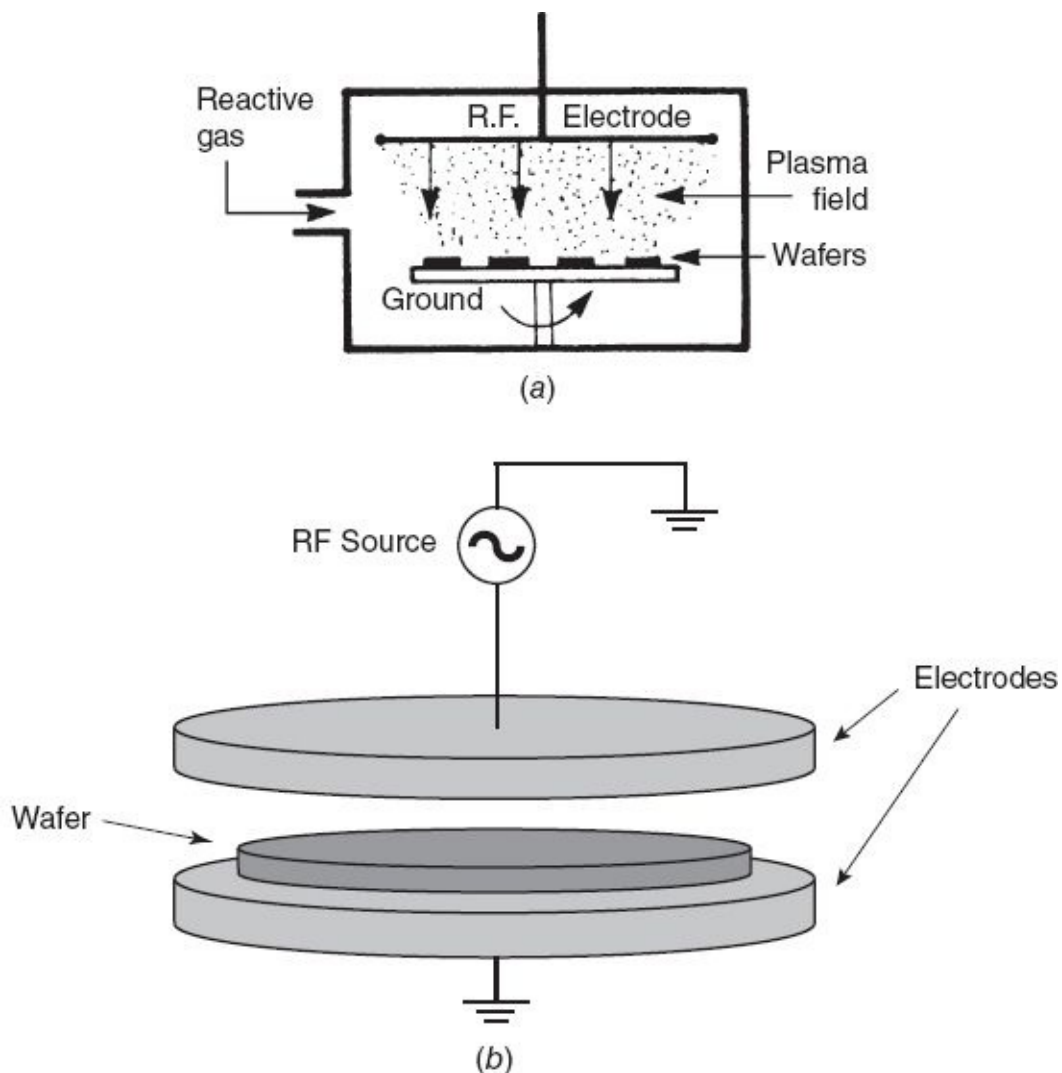


FIGURE 9.22 (a) Planar plasma etch, (b) typical single plasma etch chamber.

Planar plasma etch systems are designed in both batch and single-wafer chamber configurations. Single-wafer systems are popular for their ability to have the etch

parameters tightly controlled for uniform etching. Also, with load-lock chambers, single-wafer systems can maintain high production rates and are amenable to track based in-line automation.

RF-generated parallel plate plasma sources are giving way to new sources for 0.35- μm processing.⁸ High-density, low-pressure plasma sources under consideration are electron cyclotron resonance (ECR), high-density reflected electron, helicon wave, inductively coupled plasma (ICP), and transformer-coupled plasma (TCP).

The figures of merit for a dry-etch process are: etch rate, radiation damage, selectivity, particulate generation, post-etch corrosion, and cost of ownership.

Etch Rate

The etch rate of a plasma system is determined by a number of factors. System design and chemistry are two of them. Others are the ion density and system pressure. Ion density (no. ions/cm³) is a function of the power supplied to the electrodes. (Power supply configurations are described in [Chap. 12](#).) Increasing power creates more ions, which in turn increase the etch rate. Ion density is similar to increasing the strength of a liquid chemical etch solution. Ion densities are in the 3×10^{10} to 3×10^{12} range.⁹ System pressure influences etch rate and uniformity through a phenomenon called *mean free path*. This is the average distance a gas atom or molecule will travel before a collision with another particle. At higher pressures, there are many collisions that give the particles many directions, which in turn causes loss of edge profile control. Low pressures are preferred, but there is a trade-off with plasma damage as explained below. System pressures typically run in the 0.4 to 50 m torr range.⁹ Etch rates vary from 600 to 2000 Å/min.¹⁰

Radiation Damage

It would seem that higher-density sources with low pressure is the preferred system design. However, there is a countervailing process of *radiation* or *plasma* damage to the wafers. Within the plasma field are energetic atoms, radicals, ions, electrons, and photons.¹¹ These species, depending on their concentration and energy levels, cause various damage in semiconductors. The damage includes surface leakage, changes in electrical parameters, degradation of films (especially oxides), and damage to silicon, among others. There are two damage mechanisms. One is simple overexposure to the high-energy species in the plasma. The other is *dielectric wearout* from currents flowing across dielectrics during the etch cycle.¹¹ Higher-density sources also cause a problem for photoresist removal. The combination of the energy and low pressure tend to harden the resist to a level that is difficult to remove with conventional processes (see “Resist Stripping”). System designers are looking to plasma sources that feature high-density, low-energy ions (to reduce damage) and low-pressure operation. Besides balancing the ion density or pressure parameters, *downstream plasma* processing is an option to reduce plasma damage. The damaging species come from the high energy applied to the gas by the plasma source. Downstream systems create the plasma field in one chamber and transport it downstream to the wafer(s). The wafers are separated from the damaging plasma. To minimize the damage, the system must allow for distinguishing the plasma discharge, ionic recombination, and reduction of the electron density.¹² Downstream plasma systems were developed to minimize damage during plasma resist removal. They are attractive for etch applications even though they add more complexity to an etch system.

Selectivity

Selectivity is a major consideration in plasma-etching processes, especially when balanced against the need for overetch. Ideally, an etch time could be calculated to remove the anticipated layer thickness with just a little overetch time for safety. Unfortunately, the accumulated thickness and composition variations on multiple layer *stacks* on high-density devices present etch uniformity problems. Also on high-density devices, a phenomenon called *microloading* introduces etch-rate variations. Microloading is a change in the local etch rate relative to the area of material being etched. A large area will *load* the etching process with removed material, slowing it down in that area while a smaller etch area proceeds at a faster rate. Topography issues also drive overetch considerations. A typical situation is the opening of contact holes in both thin regions and thick regions on the device/circuit (see [Chap. 10](#)). These factors can lead to an overetch of 50 to 80 percent¹³ for metal etches and up to 200 percent for oxides and polysilicon etches.¹⁴

Overetch makes the issue of selectivity critical. There are two factors to consider: the photoresist and the underlying layer (usually silicon or silicon nitride). Dry-etch process has a higher resist-removal rate than wet processes. Given the thinner photoresist layers used to define small geometries and the increasing use of stacks of layers, the photoresist selectivity becomes critical. Compounding the selectivity issue are high-aspect-ratio patterns. Advanced devices have patterns with up to 4:1 aspect ratios or more. The holes are so narrow compared to their height that etching can slow up or stop near the bottom.¹⁵

Four methods used to control selectivity are the selection of the etching gas formula, the etch rate, the dilution of the gas near the end of the process to slow down the attack of the underlying layer, and end-point detectors in the system.

Terminating the etch process when the top layers have been removed requires an in-system end-point detector. Typically, a laser interferometer is used. A laser beam is reflected off of the wafer surface as the etch proceeds. It returns to the detector in an oscillating mode that varies with the type of material being etched. An end-point detector senses the presence of the etching layer material in the exhaust stream and automatically signals an end to the etch when no more material is detected.

Contamination, Residues, Corrosion, and Cost of Ownership

Other process problems of concern, especially in the submicron range, are particulate generation, residues, post-etch corrosion, and all the cost-of-ownership (COO) factors. One approach to reduction of particles is an electrostatic wafer holder to replace mechanical holders. Mechanical holders generate particles and cause wafer breakage, and the clamps shadow part of the wafer surface. The

electrostatic clamp holds the wafer with a direct current (dc) potential between the wafer and the chuck.¹⁶

The plasma-etch environment is highly reactive, and many chemical reactions take place. Hydroxyl groups in the photoresist react with metal halide gases to form stable metal halide (such as AlF_3 , WF_5 , WF_6) and oxides such as TiO_3 , TiO , and/or WO_2 .¹⁷ These residues create contamination problems and can interfere with the selective deposition of tungsten.¹⁸

Post-etch corrosion comes from some etch residues left on metal patterns after the etch process. The addition of copper to aluminum metal and the use of titanium or tungsten metallization increase the corrosion problem from residual chlorine after plasma etch. Minimizing this problem includes substituting fluorine-based etchants for chlorine etchants, passivating the sidewalls, and post-etch processes such as removing the residual chlorine or using a native oxide to passivate the surface.¹⁹ Other solutions include an oxygen plasma treatment, fuming nitric acid, and a wet photoresist stripper step.¹⁸ Cost-of-ownership factors are detailed in [Chap. 15](#).

[Figure 9.23](#) lists the common gas etchants for various materials. Silicon and silicon dioxide processes favor fluorine-based etchants such as CF_4 . Aluminum etching generally takes place with chlorine-based gases such as BCl_3 .

Film	Etchant	Typical Gas Compounds
Al	Chlorine	BCl_3 , CCl_4 , Cl_2 , SiCl_4
Mo	Fluorine	CF_4 , SF_4 , SF_8
Polymers	Oxygen DF_4 , SF_4 , SF_6	
Si	Chlorine, fluorine CF_4 , SF_4 , SF_6	BCl_3 , CCl_4 , Cl_2 , SiCl_4
SiO_2	Chlorine, fluorine	CF_4 , CHF_3 , C_2F_6 , C_3F_8
Ta	Fluorine	"
Ti	Chlorine, fluorine	"
W	Fluorine	"

FIGURE 9.23 Plasma etch chemicals.

Ion-Beam Etching

A second type of dry-etch system is the ion-beam system ([Fig. 9.24](#)). Unlike the chemical plasma systems, ion-beam etching is a physical process. The wafers are placed on a holder in a vacuum chamber, and a stream of argon is introduced into the chamber. Upon entering the chamber, the argon is subjected to a stream of high-energy electrons from a set of cathode (–)-anode (+) electrodes. The electrons ionize the argon atoms to a high-energy state with a positive charge. The wafers are held on a negatively grounded holder, which attracts the ionized argon atoms. As the argon atoms travel to the wafer holder, they accelerate, picking up energy. At the wafer surface, they crash into the exposed wafer layer and literally blast small amounts from the wafer surface. Scientists call this physical process *momentum transfer*. No chemical reaction takes place between the argon atoms and the wafer material. Ion-beam etching is also called *sputter etching* or *ion milling*.

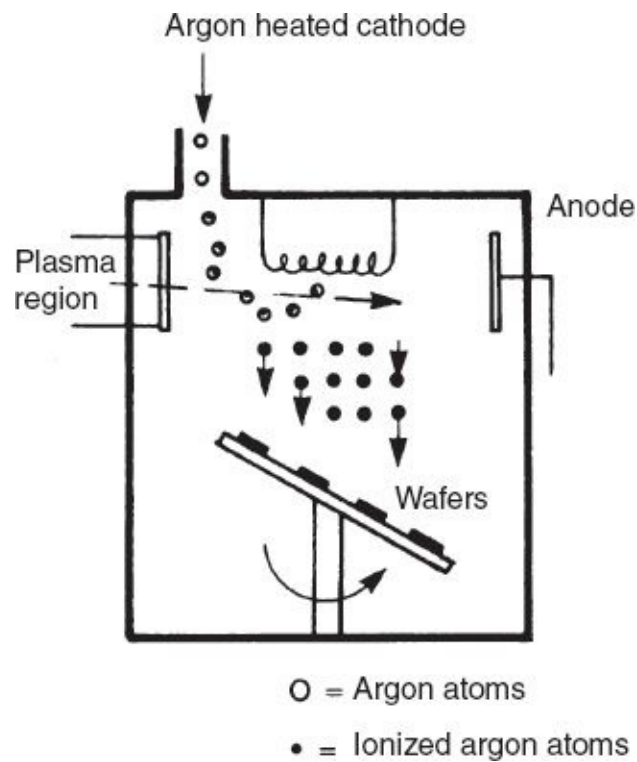


FIGURE 9.24 Ion-beam milling.

The material removal (etching) is highly directional (anisotropic), resulting in good definition of small openings. Being a physical process, ion milling has poor selectivity, especially with photoresist layers.

Reactive Ion Etching

Reactive ion etching (RIE) systems combine plasma etching and ion-beam etching principles. The systems are similar in construction to the plasma systems but have a capability of ion milling. The combination brings the benefits of chemical plasma etching along with the benefits of directional ion milling. A major advantage of RIE systems is in the etching of silicon dioxide over silicon layers. The combination etch results in a selectivity ratio of 35:1,^{[20](#)} whereas ratios of only 10:1 are available with plasma-only etching. RIE systems have become the etching system of choice for most advanced product lines.

Resist Effects in Dry Etching

For both wet-and dry-etching processes, a patterned photoresist layer is the preferred etch barrier. In wet etching, there is almost no attack of the resist by the etchants. However, in dry etching, residual oxygen in the system attacks the resist layer. The resist must remain thick enough to stand up to the etchants without becoming so thin that pinholes are present. Some structures use deposited layers as etch barriers to avoid the problem of resist lost (see [Chap. 10](#)).

Another resist-related dry-etch problem is resist baking. Within the dry-etch chamber, the temperature can rise as high as 200°C, a temperature that can bake the resist to a condition that makes it difficult to remove from the wafer. Another temperature-related problem is the tendency of resist patterns to flow and distort the images.

One unwanted effect of plasma etching is the deposition of *sidewall polymer* strings on the sides of etch patterns. The polymer comes from the photoresist. During a subsequent oxygen-plasma resist strip step, the strings become metal oxides¹⁸ that are difficult to remove.

Resist Stripping

After etching, the pattern is a permanent part of the top layer of the wafer. The resist layer that has acted as an etch barrier is no longer needed and is removed (or stripped) from the surface. Traditionally, the resist layer has been removed by wet chemical processing. Despite the issues, wet chemistry is the preferred method in the *front end of the line* (FEOL), where the surface and sensitive MOS gates are exposed and vulnerable to damage in plasma strippers.²¹ There is a growing use of plasma O₂ stripping, primarily in the *back end of the line* (BEOL), where the sensitive devices parts are covered by surface layers of dielectrics and metals.

A number of different chemicals are used for stripping. Choices depend on the wafer surface (under the photoresist), production considerations, the polarity of the resist, and the condition of the resist ([Fig. 9.25](#)). Wafers are stripped of photoresist after a number of processes: wet etch, dry etch, and ion implantation. There are different degrees of difficulty depending on the prior process. High-temperature hard bakes, plasma etch residues and sidewall polymers, and ion implantation *crusting* all present challenges for the resist removal process.

Stripper Chemistry	Strip Temperature (Centigrade)	Surface Oxide	Metallized	Resist Polarity
Acids:				
Sulfuric acid+oxidant	125	X		+/-
Organic Acids	90-110	X	X	+/-
Chromic/Sulfuric	20	X		+/-
Solvents:				
NMP/Alkanolamine	95		X	+
DMSO/Monothanolamine	95		X	+
DMAC/Diethanolamine	100		X	+
Hydroxylamine (HDA)	65		X	+

FIGURE 9.25 Wet photoresist stripper chart.

Generally, the strippers are divided into the categories of universal strippers, positive-resist-only and negative-resist-only strippers. They are also divided by the type of wafer surface: metallized or nonmetallized.

Wet stripping is used for the following reasons:

1. It has a long process history.
2. It is cost-effective.
3. It is effective in the removal of metallic ions.

4. It is a low-temperature process and does not expose the wafers to potentially damaging radiation.

Wet Chemical Stripping of Nonmetallized Surfaces

Sulfuric Acid and Oxidant Solutions

Solutions of sulfuric acid and hydrogen peroxide (SPM)²² are the most common wet strippers used for the removal of resist from *nonmetallic* surfaces. Nonmetallic surfaces are either silicon dioxide, silicon nitride, or polysilicon. This solution strips both negative and positive resists. These are the same chemical solutions and processes used for pre-tube-cleaning wafers described in [Chap. 7](#).

Nitric acid is sometimes used as an additive oxidant in a sulfuric acid bath. A mixing ratio of about 10:1 is typical. A drawback to nitric acid is that it turns the bath a light orange color, which can mask the buildup of carbon in the bath. All of these solutions dissolve the resist by an oxidation mechanism.

Wet Chemical Stripping of Metallized Surfaces

Stripping from metallized surfaces is a more difficult task, because the metals are subject to attack or oxidation. Four types of wet chemicals are used for stripping metallized surfaces. They are:

1. Organic strippers
2. Solvent strippers
3. Solvent-amine strippers
4. Specialty strippers

Phenolic Organic Strippers

Organic strippers contain a combination of sulfonic acid (an organic acid) and chlorinated hydrocarbon solvents such as duodexabenzene. The formula requires phenol to create a rinsable solution. In the 1970s, concern over the toxic ingredients in these formulas led to the development of sulfonic acid, nonphenolic, nonchlorinated²³ resist strippers. Stripping the photoresist requires heating the solutions into the 90° to 120°C range. Often, the process uses two or three heated strip baths. Rinsing is in two steps, the first being a solvent followed by a water rinse and drying step.

Solvent-Amine Strippers

One of the advantages of positive resists is their ease of removal from the wafer surface. A positive-resist layer that has not been hard baked is easily removed from the wafer with a simple acetone soak. In fact, acetone has been the traditional positive-resist stripper. Unfortunately, acetone represents a fire hazard, and its use is discouraged.

Several manufacturers supply positive-only strippers based on solvent and organic amine solutions. *N*-methyl pyrrolidine (NMP)²⁴ is the most used solvent. Others are dimethylsulfoxide (DMSO), sulfolane, dimethylformamide (DMF), and dimethylacetamide (DMAC). These strippers are effective, water-rinsable, and drain-dumpable. The strippers may be heated to increase the removal rate and/or to remove resist films that have been through a high-temperature hard bake. Solvent and solvent-amine strippers remove the resist by a chemical dissolution mechanism.

Specialty Wet Strippers

A number of wet chemical strippers have been developed to solve specific problems. One is a positive-resist stripper based on hydroxylamine (HDA) chemistry.¹⁸ Another chemistry relies on chelating agents to, in effect, chemically bind metal contaminants in the solution.²⁵ The stripper removes a host of plasma etch residues and polyimide layers not removed by solvent-amine strippers. Other strippers include corrosion inhibitors.²⁶

[Figure 9.25](#) is a table of the most common wet resist strippers and their uses. The advent of multilayer metallization systems with transition metal-connecting plugs requires wet strippers that do not attack these metals.

Dry Stripping

Like etching, the dry-plasma process can also be applied to resist stripping. The wafers are placed in a chamber, and oxygen is introduced ([Fig. 9.26](#)). The plasma field energizes the oxygen to a high-energy state, which in turn oxidizes the resist components to gases that are removed from the chamber by the vacuum pump. The term *ashing* is used to designate plasma processes designed to remove only organic residues. Plasma stripping indicates a process designed to remove both organic and inorganic residues. In dry strippers, the plasma is generated by microwave, RF, and UV-ozone sources.²⁷

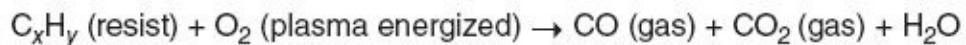


FIGURE 9.26 Resist removal by plasma oxygen.

The major advantage of plasma resist stripping is the elimination of wet hoods and the handling of chemicals. The principal disadvantage is its ineffectiveness in the removal of metallic ions. There is not enough energy in the plasma field to volatilize the metallic ions. Another consideration of plasma stripping is radiation damage to the circuits from the high-energy plasma field. This problem is reduced with system designs that have the plasma chamber removed from the stripper chamber. They are called *downstream strippers*, because the plasma is created downstream from the wafers. MOS wafers are more sensitive to radiation effects during stripping.

Replacement of wet stripping by dry or plasma techniques has been a long-term industry project. However, the inability of oxygen plasma to remove mobile ionic metal contamination and certain metal residues, and radiation damage concerns, have maintained wet stripping and wet-dry combinations as the mainstream photoresist removal processes. Plasma stripping is used to remove hardened resist layers. A following wet strip step is used to remove the residuals not removed by the plasma.

Post-Ion Implant and Plasma Etch Stripping

Two problem areas of dry-resist stripping are after ion implant and after a plasma strip. Ion implant causes extreme polymerization of the resist and crusting of the top. Generally, the resist is removed or reduced with a dry process followed by a wet process. Post-plasma etch-resist layers are similarly difficult to remove. In

addition, the etch process can leave residues, such as AlCl_3 and/or AlBr_3 , that react with water or air to form compounds that corrode the metal system.²⁸ Low-temperature plasma can remove the offending compounds before they take on a corrosive chemistry. Another approach is to add halogens to the plasma atmosphere to minimize the formation of the insoluble metal oxides. This is another instance of setting process parameters to achieve efficient processing (resist removal) without inducing wafer-surface damage or metal corrosion.

New Stripping Challenges

The stripping processes and chemistry described are traditional and fairly simple from a technical perspective. Development of ever-smaller dimensions, larger and more dense chips, III-V and SiGe substrates, shallower junctions, multilayer stacks with deep vias, copper dual damascene processes, and others have driven changes to the resist-stripping processes.

Resist-stripping effectiveness is very dependent of its history through the expose, developing, and bake steps. Collectively they present different challenges to resist stripping. Hence the traditional chemistries have evolved into specialized formulas. There are stripping solutions for plasma-hardened resists, etch residues, in conjunction with low-k technologies, and so on. New device structures with shallower junctions and narrow gate widths require resist stripping processes that do not etch the exposed wafer surface or leave electrically active residues.²⁸

Final Inspection

The final step in the basic photomasking process is a surface inspection. It is essentially the same procedure as the develop inspect, with the exception that the majority of the rejects are fatal (no rework is possible).

The one exception is contaminated wafers that may be recleaned and reinspected. Final inspection certifies the quality of the outgoing wafers and serves as a check on the effectiveness of the develop inspection. Wafers that should have been identified and pulled from the batch at develop inspect are rejected from the batch.

The wafers receive a first surface inspection in incident white or ultraviolet light for stains and large particulate contamination. This inspection is followed by a microscopic or automatic inspection for defects and pattern distortions. Measurement of the critical dimensions for the particular mask level is also part of the final inspection. Of primary interest is the quality of the etched pattern, with underetching and undercutting being two parameters of concern.

Mask Making

In [Chap. 5](#), the steps of circuit design were detailed. In this section, the process used to construct a photomask or reticle is examined. Originally, the masks were made from emulsion-coated glass plates. The emulsions are similar to those found on camera film. These masks were vulnerable to scratches, deteriorated during use, and were not capable of resolving images in the sub-3- μm range. Masks for most modern work use a chrome-on-glass technology. This mask-making technology is almost identical to the basic wafer-patterning operation ([Fig. 9.27](#)).

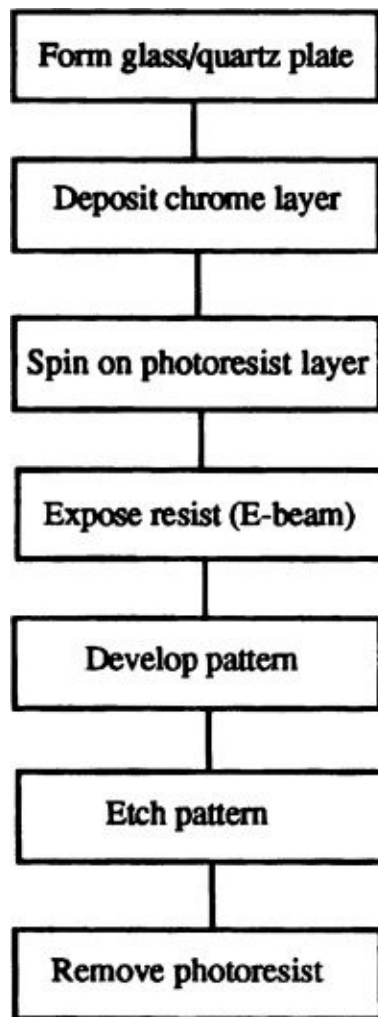


FIGURE 9.27 Major steps in mask/reticle plate processing.

In fact, the goal is the same—the creation of a pattern in the thin chrome layer on the glass reticle surface. The preferred materials for mask or reticles are borosilicate glass or quartz, which have good dimensional stability and

transmission properties for the wavelengths of the exposing sources. Chrome layers are in the 1000-Å range and deposited on the glass by sputtering (see [Chap. 12](#)). Advanced mask or reticles use layers of chromium, chromium oxide, and chromium nitride.²⁹

Chrome layers are effective energy blockers at wavelengths of 365 nm, 248 nm, and 193 nm. Smaller dimensions require different exposure sources (EUV, X-ray, electrons, and ions), which in turn these require entirely new materials for the substrate and the pattern film. (See [Chap. 10](#).)

Mask or reticle making follows a number of different paths depending on the starting exposure method (pattern generation, laser, e-beam) and the end result (reticle or mask) ([Fig. 9.28](#)). Flow A shows the process for making a reticle using a pattern generator, which is an older technology. A pattern generator consists of a light source and a series of motor-driven shutters. The chrome-covered mask or reticle, with a layer of photoresist, is moved under the light source as the shutters are moved and opened to allow precisely shaped patterns of light to shine onto the resist, creating the desired pattern. The reticle pattern is transferred to the resist-covered mask blank by a step-and-repeat process to create a master plate. The master plate is used to create multiple working mask plates in a contact printer. This tool brings the master into contact with a resist-covered mask blank and has a UV light source for transferring the image. After each of the exposure steps (pattern generation, laser, e-beam, master plate expose, and contact print), the reticle or mask is processed through development, inspection, etch, strip, and inspection steps that transfer the pattern permanently into the chrome layer. Inspections are very critical, since any undetected mistake or defect has the potential of creating thousands of scrap wafers. Reticles for this use are generally 5 to 20 times the final image size on the mask.³⁰

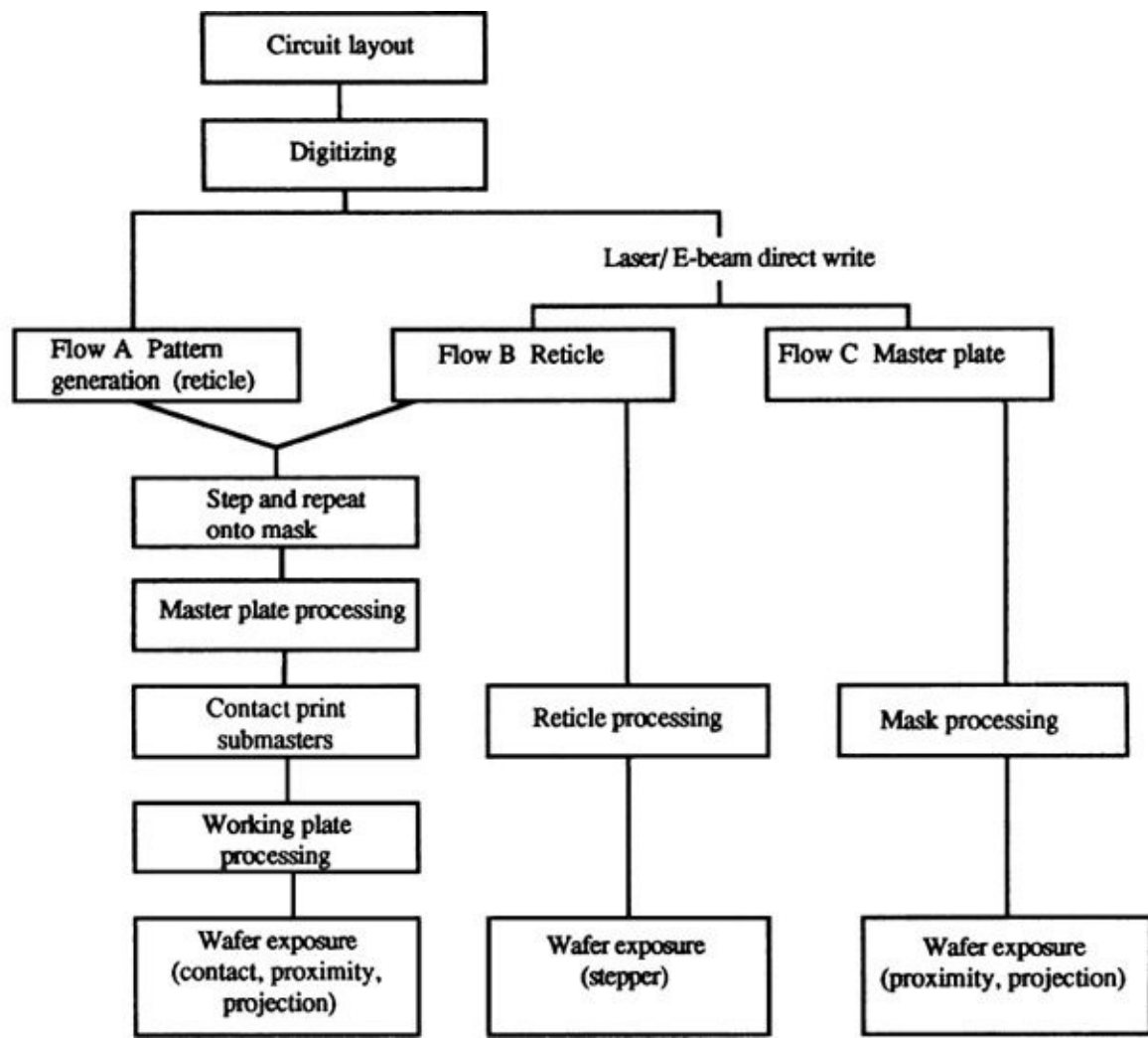


FIGURE 9.28 Mask-or reticle-making processing flows.

Advanced products with very small geometries and tight alignment budgets require high-quality reticle and/or masks. The reticles and masks for these processes are made with lasers or e-beam direct write exposure (flows A&B). Laser exposure uses a wavelength of 364 nm, making it an I-line system. It allows using standard optical resists and is faster than the e-beam. Direct-write laser sources are turned on and off with an acoustooptical modulator (AOM).³¹ In all cases, the reticle or mask is processed to etch the pattern in the chrome. Other mask or reticle process flows may be employed. The reticle in flow A may be laser or e-beam generated, or the master plate may be laser-or e-beam generated.

VLSI and ULSI-level circuits require virtually defect-free and dimensionally perfect masks and reticles. Critical dimension (CD) budgets from all sources are 10 percent or better, leaving the reticles with a 4 percent error margin.³² There are procedures to eliminate unwanted chrome spots and pattern protrusions with laser “zapping” techniques. Focused ion beams (FIBs) is the preferred repair technology for small image masks and reticles. Clear or missing pattern parts are “patched”

with a carbon deposit. Opaque or unwanted chrome areas are removed by sputtering from the beam.

Summary

For VLSI and ULSI work, the resolution and registration requirements are very stringent. In 1977, the minimum feature size was 3 μm . By the mid-1980s, it had passed the 1- μm barrier. By the 1990s, 0.5- μm sizes were common with 0.35- μm technology planned for production circuits. Circuit design projections call for minimum gate sizes of 10 to 15 nm in 2016.³³

Chip manufacturers calculate several *budgets* for each circuit product. A *critical dimension (CD) budget* calculates the allowable variation in the image dimensions on the wafer surface. For products with submicron minimum feature sizes, the CD tolerances are 10 to 15 percent.³⁴ Also of concern is the critical defect size relative to the minimum feature size. These two parameters are brought together in an *error budget* calculated for the product. An *overlay budget* is the allowable accumulated alignment error for the entire mask set. A rule of thumb is that circuits with micron or submicron feature sizes must meet registration tolerances of one-third the minimum feature. For a 0.35- μm product, the allowable overlay budget is about 0.1 μm .³⁵

Review Topics

Upon completion of this chapter, you should be able to:

1. Draw a cross-section of a wafer before and after developing.
 2. Make a list of the developing methods.
 3. Explain the purpose and methods of hard bake.
 4. List at least five reasons why a wafer can be rejected at develop inspect.
 5. Draw a diagram of the develop-inspect-rework loop.
 6. Explain the methods and relative merits of wet and dry etch.
 7. Make a list of the resist strippers used to strip photoresist from oxide and metal films.
 8. Explain the purpose and methods of final inspection.
-

References

- [1.](#) Elliott, D., *Integrated Circuit Fabrication Technology*, 1976, McGraw-Hill, New York, NY:216.
- [2.](#) Busnaina, A., and Dai, F., "Megasonic Cleaning," *Semiconductor International*, Aug. 1997.
- [3.](#) Wolf, S., and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Newport Beach, CA:530.
- [4.](#) Singer, P., "Meeting Oxide, Poly and Metal Etch Requirements," *Semiconductor International*, Cahners Publishing, Apr. 1993:51.
- [5.](#) Ibid., p. 51.
- [6.](#) Wolf, S., and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Newport Beach, CA:532.
- [7.](#) Murray, C., "Wet Etching Update," *Semiconductor International*, May 1986:82.
- [8.](#) Burggraaf, P., "Advanced Plasma Sources: What's Working?" *Semiconductor International*, Cahners Publishing, May 1994:57.
- [9.](#) Singer, P., "Meeting Oxide, Poly and Metal Etch Requirements," *Semiconductor International*, Cahners Publishing, April 1993:53.
- [10.](#) Elliott, D., *Integrated Circuit Fabrication Technology*, 1976, McGraw-Hill, New York, NY:275.
- [11.](#) Fonsh, S., Viswanathan, C., and Chan, Y., "A Survey of Damage Effects in Plasma Etching," *Solid State Technology*, PennWell Publishing Company, Jul. 1994:99.
- [12.](#) Boitnott, C., "Downstream Plasma Processing: Considerations for Selective Etch and Other Processes," *Solid State Technology*, PennWell Publishing Company, Oct. 1994:51.
- [13.](#) Riley, P., Pengm, S., and Fang, L., "Plasma Etching of Aluminum for ULFI Circuits," *Solid State Technology*, PennWell Publishing Company, Feb. 1993:4.
- [14.](#) Engelhardt, M., "Advanced Polysilicon Etching in a Magnetically Confined Reactor," *Solid State Technology*, PennWell Publishing Company, Jun. 1993:57.
- [15.](#) Singer, P., "Meeting Oxide, Poly and Metal Etch Requirements," *Semiconductor International*, Cahners Publishing, Apr. 1993:51.
- [16.](#) Newboe, B., "Wafer Chucks Now Have an Electrostatic Hold," *Semiconductor International*, Cahners Publishing, Feb. 1993:30.

- [17.](#) Cardinaud, C., Peignon, M., and Turban, G., "Surface Modification of Positive Photoresist Mask during Reactive Ion Etching of Si and W in SF₆ Plasma," *J. Electro-chemical Soc.*, vol. 198, 1991:284.
- [18.](#) Lee, W. M., *A Proven SubMicron Photoresist Stripper Solution for Post Metal and Via Hole Processes*, 1993, EKC Technology, Inc., Hayward, CA.
- [19.](#) Clayton, F., and Beeson, S., "High-Rate Anisotropic Etching of Aluminum on a Single-Wafer Reactive Ion Etcher," *Solid State Technology*, PennWell Publishing Company, Jul. 1993:93.
- [20.](#) Elliott, D., *Integrated Circuit Fabrication Technology*, 1976, McGraw-Hill, New York, NY:282.
- [21.](#) Dejule, R., "Managing Etch and Implant Residue," *Semiconductor International*, Aug. 1997:62.
- [22.](#) EKC Technology Inc., Technical Bulletin SA-80, 1999.
- [23.](#) EKC Technology Inc., Technical Bulletin—Nophenol 922, 1999.
- [24.](#) EKC Technology Inc., Technical Bulletin—Posistrip Series, 1999.
- [25.](#) Dejule, R., "Managing Etch and Implant Residue," *Semiconductor International*, Aug. 1997:57.
- [26.](#) Levenson, M. D., "Wet Stripper Companies Clean Up," *Solid State Technology*, PennWell Publishing Company, Apr. 1994:31.
- [27.](#) Burggraaf, P., "What's Driving Resist Dry Stripping?" *Solid State Technology*, PennWell Publishing Company, Nov. 1994:61.
- [28.](#) Berry III, I. L., Waldfried, C., Roh, D., et al., *Photoresist Strip Challenges for Advanced Lithography at 20nm*, www.axcelis.com.
- [29.](#) Dejule, R., "Managing Etch and Implant Residue," *Semiconductor International*, Aug. 1997:58.
- [30.](#) Grenon, B., "A Comparison of Commercially Available Chromium-Coated Quartz Mask Substrates," *OCG Microlithography Seminar*, Interface 94:37.
- [31.](#) Wolf, S., and Tauber, R. N., *Silicon Processing*, vol. 1, Lattice Press, 2000, Sunset Beach, CA:477.
- [32.](#) Reynolds, J., "Mask Making Tour Video Course," *Semiconductor Services*, Redwood City, CA, Aug. 1991, Segment 5.
- [33.](#) Reynolds, J., "Elusive Mask Defects: Random Reticle CD Variation," *Solid State Technology*, PennWell Publishing Company, Sep. 1994:99.
- [34.](#) Semiconductor Industry Association, *International Technology Roadmap for Semiconductors*, 2001, 2002, www.semiconductors.org/.
- [35.](#) Wiley, J., and Reynolds, J., "Device Yield and Reliability by Specification of Mask Defects," *Solid State Technology*, PennWell Publishing Company, Jul.

1993:65.

36. Simon, K., "Abstract-Alignment Accuracy Improvement by Consideration of Wafer Processing Impacts," *SPIE Symposium on Microlithography*, 1994:35.

CHAPTER 10

Next Generation Lithography

Introduction

Feature sizes decreasing to the nanometer range, the increasing need for low defect densities, increasing chip density and size, along with larger-diameter wafers have challenged the industry to squeeze every capability out of traditional processes and develop new ones. The problems encountered and current solutions to reaching nanometer circuit dimensions are explored in this chapter. These processes and technologies are known collectively as *next-generation lithography* (NGL).

In this chapter, some of the limits of optical imaging and advanced process solutions are presented. Also, industry developments extending optical lithography and development of NGL have proceeded on almost all of the individual elements of basic patterning processes. They include resist development, mask materials and designs, exposure sources alignment and exposure schemes, reflection control, and process schemes. In the present and future eras, lithography advances along the technology node scale will involve combinations of these factors. No single lithography process step is expected to provide a comprehensive breakthrough.

Challenges of Next Generation Lithography

The ten-step patterning process detailed in [Chaps. 8](#) and [9](#) is a one photoresist-layer basic process. It would be sufficient for the production of medium scale integration (MSI) and some simple large-scale integration (LSI) and very large-scale integration (VLSI) circuits. However, the VLSI or ultra large-scale integration (ULSI) demands of decreasing feature sizes and defect densities is beyond the capabilities of these basic processes. The limits showed up at the 2–3 μm level and became critical in the submicron era. Problems included physical limitations associated with the optical exposure equipment, resolution limits of photoresists, and a host of surface problems, including reflective surfaces and multilevel topographies.

In the mid-1970s, it was widely accepted that optical photoresist processes had a lower resolution limit of about 1.5 μm . This projection gave rise to the interest in X-ray and e-beam exposure systems.

However, manipulation and improvements on the basic processes have successively lowered the usable range of optical lithography to the 0.2- μm range.¹ In the first edition of *Microchip Fabrication* (1984), I reported that “industry futurists project that either e-beam or X-ray exposure will replace the UV and DUV sources by the mid-1990s.” It did not happen. Optical lithography has been walking the plank for decades. And every generation of engineers has tweaked and improved patterning based on optical exposure systems to bring the industry to the 100-nm node.² The crystal ball of the past has been replaced with the SIA’s *International Technology Roadmap for Semiconductors* (ITRS). [Figure 10.1](#) shows the various “nodes” of future devices and their year of introduction.³ A node is essentially the gate width. Lithography requirements are also by the half-pitch of as characterized as the 1/2 pitch width of adjacent lines and space ([Fig. 10.2](#)).

Years of Production		2012	2013	2015	2016	2018
Technology Node (hp)		hp45	hp32		hp22	
Dram 1/2 Pitch (hp)	nm	35	32	25	22	18
MPU Printed Gate Width	nm	20	18	14	13	10
MPU Physical Gate Width	nm	14	13	10	9	7

FIGURE 10.1 ITRS technology node projection.

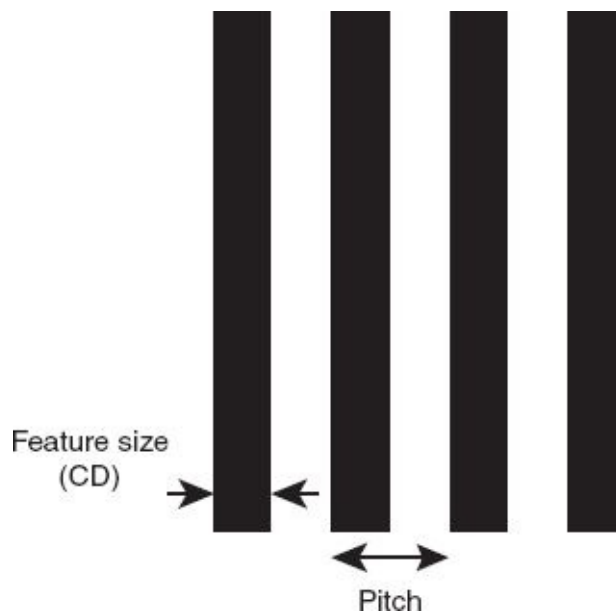


FIGURE 10.2 Feature size and pitch.

There is a basic relationship between the components of an exposure system. The relationship is:

$$\sigma = k \frac{\lambda}{\text{NA}} \quad (10.1)$$

where σ = the minimum feature size

k = constant (sometimes called the Rayleigh constant)

λ = the wavelength of the exposing light

NA = is the numerical aperture of the lens

The term k (or k_1 in some formulas) relates to the ability of the lens (or total lithography system) to differentiate two adjacent images. Diffraction effects will cause even the most perfect lens to blur images as they get closer to each other. Values of k are about 0.5.

The formula shows that decreasing the wavelength and/or increasing the NA are two ways to print a smaller image size. There are other factors also used in NGL. All are addressed using the schematic of the basic components of an NGL lithographic tool ([Fig 10.3](#)).

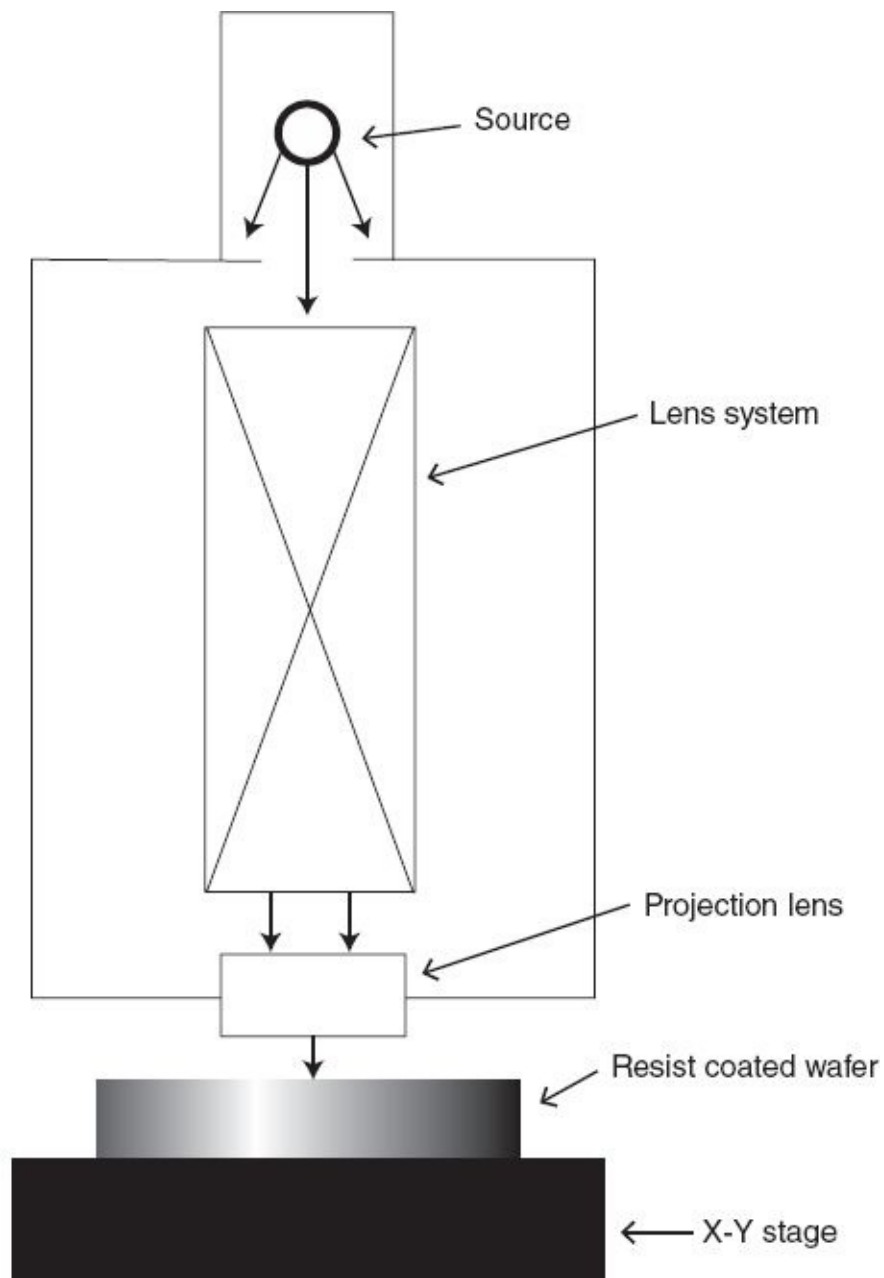


FIGURE 10.3 Lithographic tool system components.

The Rayleigh formula clearly shows that using a smaller-wavelength exposure source increases the ability to pattern smaller features. The industry has done exactly that over the years. [Figure 10.4](#) shows the progression in smaller feature requirements and the accompanying exposure sources.

Technology Generation	Wavelength (nm)	Exposure Source
DUV G line	436	Mercury (Hg)
DUV H line	405	"
DUV I line	365	"
Excimer Laser	248	KrF
"	193	ArF
"	157	F ₂
Extreme UV (EU)	135	Tin Vapor

FIGURE 10.4 Exposure sources.

High-Pressure Mercury Lamp Sources

Exposure sources are chosen to match the spectral response characteristic of the resist and the feature size of the images. Early optical aligner systems used a high-pressure mercury lamp, which produces light as a current is passed through a tube of mercury. The high-pressure atmosphere allows a higher level of stimulation of the mercury without evaporation.

The energy from the mercury comes out in bundles of waves grouped in ranges. Some resists are designed to react to the entire range of the wavelengths and some to specific wavelengths. Other resists are designed to respond to the specific high-energy peaks of a mercury lamp. The three peaks are at the 365-, 405-, and 436-nm wavelengths.

They are referred to as G, H, and I lines, respectively. Steppers used for advanced imaging often have filters to expose the resist to only the G-or I-lines. I-line exposure has been the source of choice in the submicron era. The I-line has a wavelength of 365 nm (or 0.365 μm), which is near the image size for 0.350 products. The problem is the physical difficulty of resolving images smaller than the wavelength of the light.

Excimer Lasers

There are also DUV sources available from excimer lasers. The following gas lasers and their associated wavelength are used: XeF (351 nm), XeCl (308 nm), KrF (248 nm), and ArF (193 nm).⁴ Both KrF and ArF have been developed for resolving the 130-nm node processes, while ArF is considered the source of choice for below 130-nm patterning.⁵ F₂ is also considered a candidate to follow ArF.⁶

Extreme Ultraviolet

Extreme ultraviolet (EUV) is next in the parade of exposure sources with smaller wavelengths. With a wavelength of 13.5 nm, images in the 18-to 24-nm range are possible. A tin vapor, turned into a plasma, is the source element. ASML uses two approaches. In one scheme called laser-produced plasma (LPP), a droplet of molten tin produces EUV light from a high-energy laser blast. In another called laser-assisted discharge plasma (LDP) an electrical charge is put through the tin vapor to produce EUV photons.⁷ Since glass absorbs EUV photons, the exposure system uses extremely flat mirrors to direct the beam. And since air also absorbs the photons, the entire process has to take place in a vacuum. However, system production rate is not met through production requirements. Hence, initial use will be on the critical layers in a mix-and-match arrangement with immersion lithography tools.

X-Rays

The desire for higher-resolution exposure sources inevitably led to the consideration of two nonoptical sources: X-ray sources and electron beams (e-beams). X-rays are high-energy photons with wavelengths of 4 to 50 Å.⁸ This range of wavelengths is capable of very small image sizes (down to the 0.1-μm level) due to the lack of diffraction effects. X-ray aligners are projection systems ([Fig. 10.5](#)) using a full-size mask (1:1 mask to wafer image). They generally have higher output through shorter exposure times. Reflection and scattering in the resist is minimal, and there are few depth-of-focus problems. X-ray-exposed wafers show a lower level of defects from dust and organic matter on the mask, because the X-rays pass through the spots.

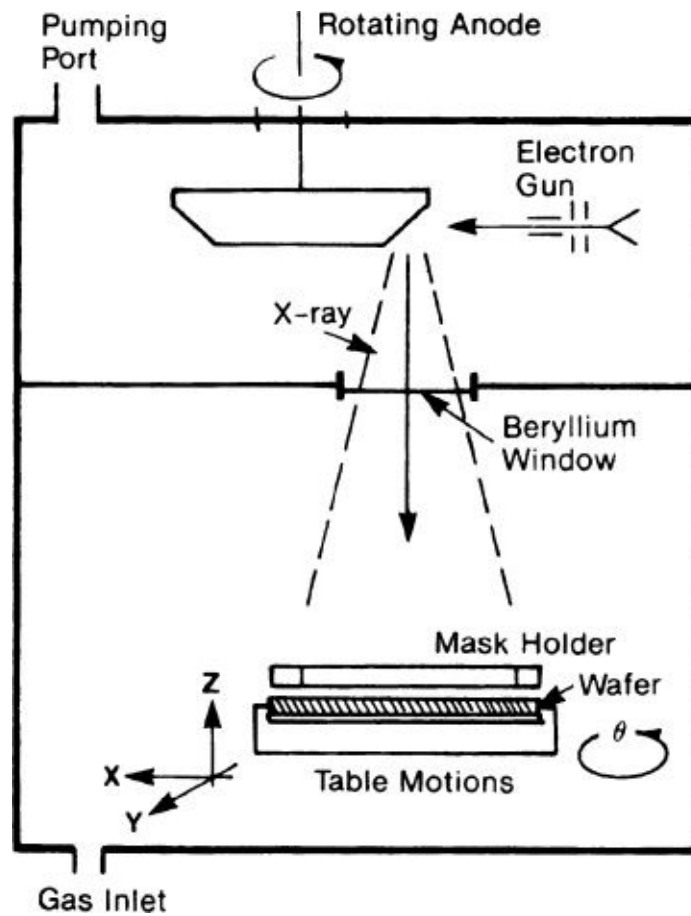


FIGURE 10.5 X-ray exposure system.

A number of difficulties have surfaced with production-level X-ray aligners. One is the development of X-ray-blocking masks. Because X-rays pass through conventional chrome and glass masks, a process that requires gold as the blocking layer and other materials that stand up to the high energy of the X-rays needs to be developed (see “Maskmaking”).

While development work goes on in defining X-ray equipment, it also goes on in developing X-ray resists. This work is complicated, because there is no standard X-ray source, and the resists must show high sensitivity to X-rays as well as being good etch barriers. These last two factors have proven difficult to balance in resist chemistries. Another barrier limits X-ray aligners to a 1:1 printing. The high-energy X-rays destroy conventional optics needed for reduction systems. And the limitation of reticle sizes means that only steppers are practical for production machines.

X-ray sources include standard X-ray tubes or laser-driven sources as point sources or synchrotron generators. Point sources are similar to conventional systems with one exposure source per machine. Synchrotrons are large, expensive machines that can accelerate electrons in an orbit. The orbiting electrons give off X-rays in a process called *synchrotron radiation*. As the X-rays rotate, they can be

directed through ports to a number of individual aligners.

Electron Beam or Direct Writing

Electron-beam lithography is a mature technology used in the production of high-quality masks and reticles. The system (Fig. 10.6) consists of an electron source that produces a small-diameter spot and a “blanker” capable of turning the beam on and off.

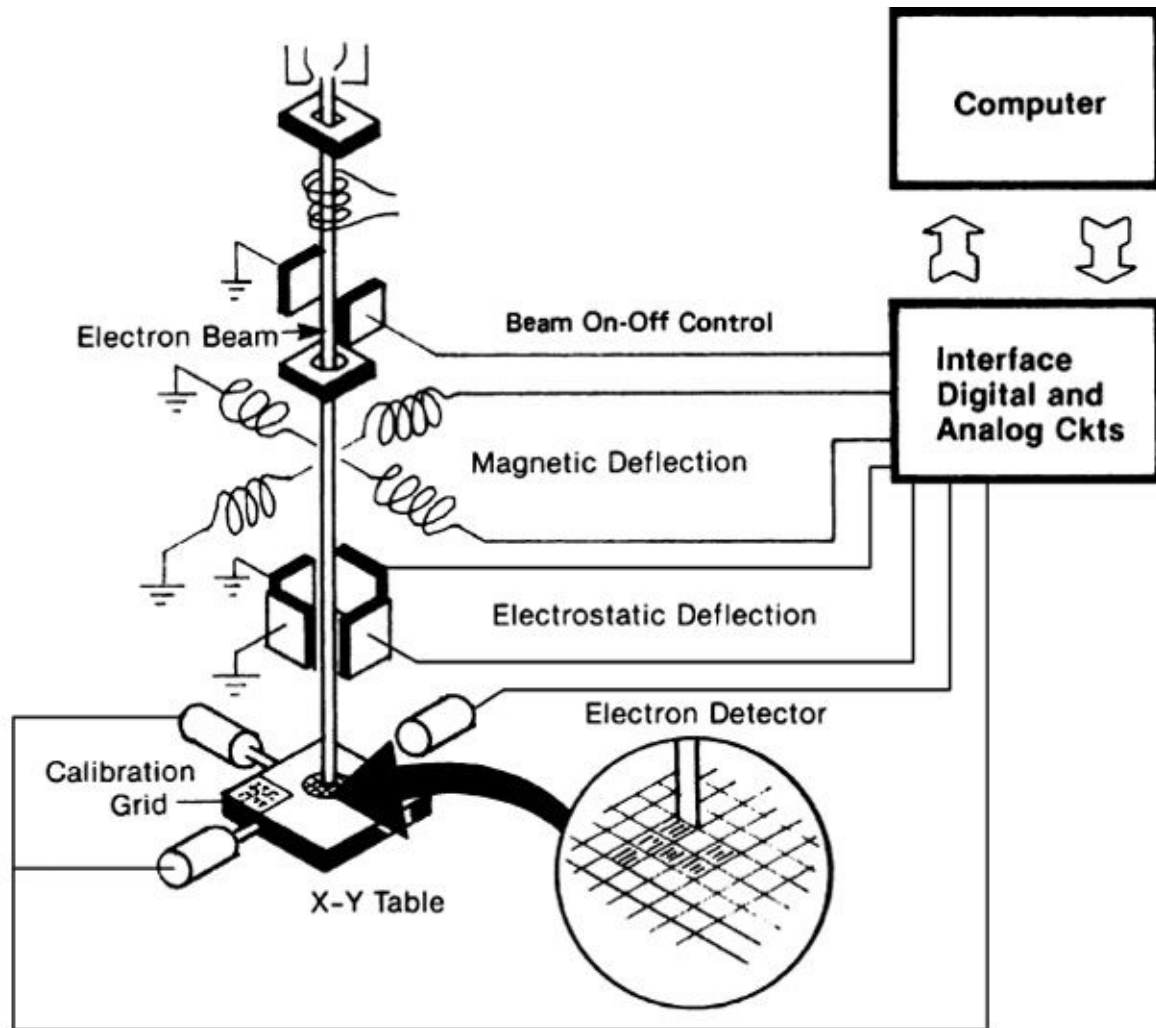


FIGURE 10.6 Electron-beam exposure system.

The exposure must take place in a vacuum to prevent air molecules from interfering with the electron beam. The beam passes through electrostatic plates capable of directing (or steering) the beam in the x-y direction onto the mask, reticle, or wafer. This system is functionally similar to the beam-steering mechanisms of a television set. Precise direction of the beam requires that the beam travel in a vacuum chamber in which there is the electron beam source, support

mechanisms, and the substrate being exposed.

There is no mask or reticle used to generate the pattern. With no mask, a source of defects and errors is eliminated along with the expense of the mask or reticle. The blanking (beam on and off) and steering functions are controlled by a computer that has in its memory the wafer pattern taken directly from the computer-aided design (CAD) stage. The beam is directed to specific positions on the surface by the deflection subsystem and the beam turned on where the resist is to be exposed. Larger wafers are mounted on an x-y stage and are moved under the beam to achieve full surface exposure. This alignment and exposure technique is called *direct writing*.

The pattern is exposed in the resist by either raster or vector scanning ([Fig. 10.7](#)). Raster scanning is the movement of the electron beam side to side and down the wafer. The computer directs the movement and activates the blanker in the regions where the resist is to be exposed. One drawback to raster scanning is the time required for the beam to scan, since it must travel over the entire surface. In vector scanning, the beam is moved directly to the regions that have to be exposed. At each location, small square-or rectangular-shaped areas are exposed, until the desired shape is exposed.

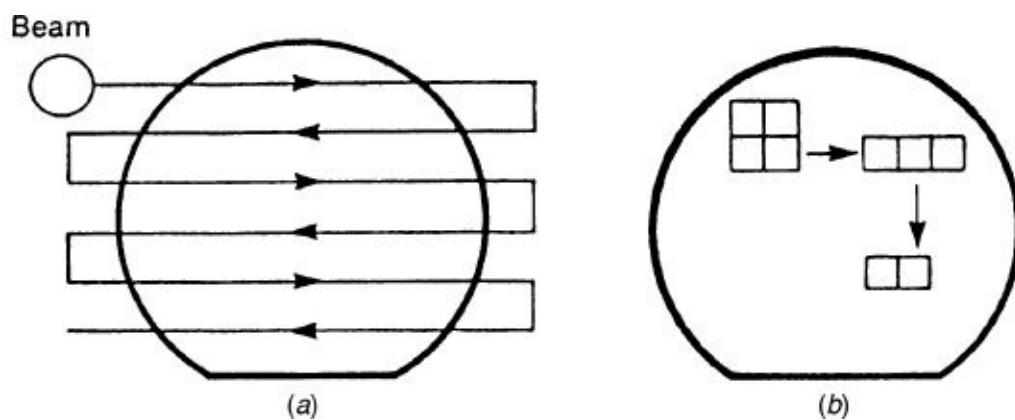


FIGURE 10.7 Electron-beam scanning.

Alignment and overlay parameters are very good with electron-beam systems, because no distortions are introduced from masks or from optical effects such as diffraction. Resolution is also good, with current machines capable of 0.25- μm feature sizes.⁹ Drawbacks to full use of electron-beam systems in wafer production are speed and cost. A factor in the slowness of the system is the time required to create the vacuum and release it in the exposing chamber.

The basics of an electron beam system have been described. Unfortunately, mask-less (raster or vector scanning) systems are too slow for patterning onto wafer surfaces. An advanced system using e-beam exposure is electron-beam projection lithography (EPL). This system uses an e-beam exposure source, a mask, and a

scanning projection method. A system developed by Lucent Technologies, called SCALPEL, uses a scattering mask ([Fig. 10.8](#)). Electrons are highly energetic and pass through most materials. While conventional masks block portions of the exposing beam and allow other portions to pass through, a scattering mask allows passage of the e-beam through both segments. However, one part of the mask scatters the e-beam as shown. A reduction lens focuses the beam down toward the wafer. In between is an aperture that essentially allows the unimpeded beam to pass onto the wafer surface and blocks the scattered beams.

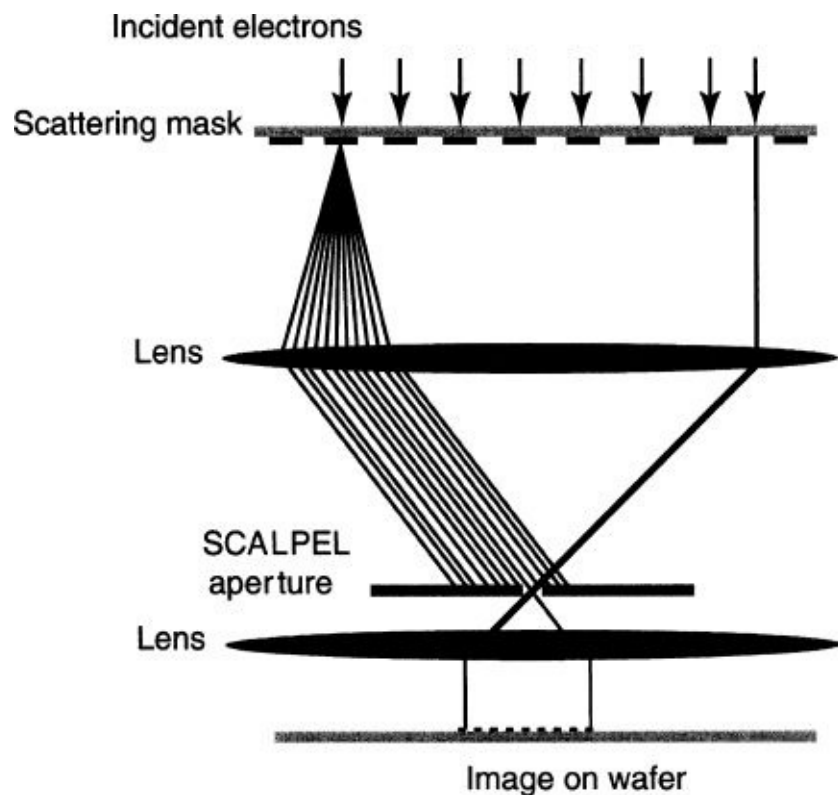


FIGURE 10.8 Basic SCALPEL principle of operation showing contrast generation by differentiating more or less scattered electrons. (*Bell Laboratories, Lucent Technologies Website.*)

Current mask design for this system uses a multilayer structure. The “pattern” is defined in a layer of silicon nitride that was deposited on a silicon wafer. Most of the wafer is etched away, except for silicon “struts” left to keep the mask rigid. This arrangement is mounted on a thin membrane. All of the parts of this mask are transparent to the e-beam. However, the pattern is defined because the incoming beam is broken into two components: scattered and nonscattered. Because this pattern is scanned onto the wafer surface, the struts do not come into the scanned pattern.

Numerical Aperture of a Lens

Early semiconductor imaging used contact or proximity exposure systems where the wafer and mask are touching or in close proximity. But small feature products are exposed with *projection systems* where the mask/reticle and wafer are separated. Projection optics presents particular problems. The challenge is to project the image from the mask or reticle to the wafer surface with as little loss of resolution or dimensional control as possible. Small image sizes require the use of short wavelengths as addressed above. Projection systems use a lens to focus the exposure beam onto the resist or wafer surface.

The minimum image that can be resolved on the wafer is constrained by the physical attributes of the projection optical system. A factor in a lens system is the numerical aperture (NA). The NA is the ability of the lens to gather light. The relationship is also shown previously in the Rayleigh formula (Eq. 10.1).

The term k (or k_i in some formulas) relates to the ability of the lens (or total lithography system) to differentiate two adjacent images. Diffraction effects will cause even the most perfect lens to blur images as they get closer to each other. Values of k are about 0.5.

The formula shows that decreasing the wavelength, as described above, is one solution to printing smaller sizes. Increasing the lens NA is website. [Figure 10.9](#) shows the basic lens focus dynamic. However, there are limits to increasing the NA. That is because of a trade-off parameter called *depth of focus* or *depth of field* (DoF). In regular photography, we run into depth-of-field problems when the foreground is in focus and the background is out of focus, or vice versa. The relationship of the lens parameters for these factors is shown in [Fig. 10.10](#).

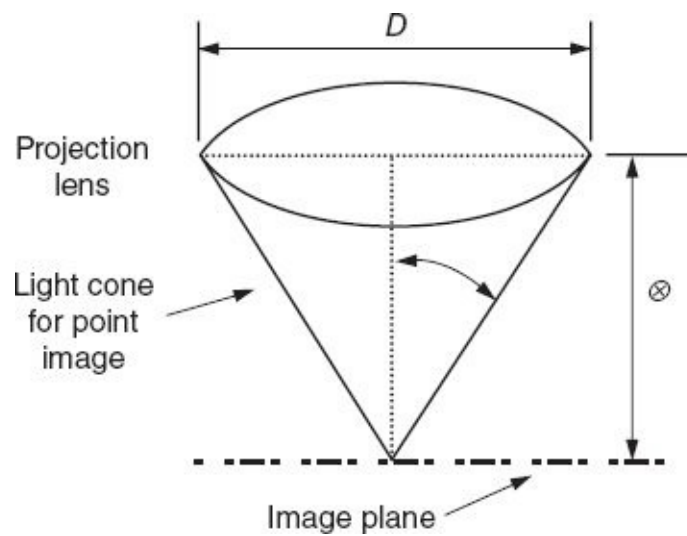


FIGURE 10.9 Lens focus relationships.

Lens examples Wavelength	NA	k_1	Resolution (um)	DoF (um)
I-line	0.62	0.48	0.28	0.95
KrF	0.82	0.36	0.11	0.37
ArF	0.92	0.31	0.065	0.23
F ₂	0.85	0.31	0.057	0.22

FIGURE 10.10 Source and lens factors.

Images, both on the top surface and at the bottom of valleys, must have good resolution and the correct dimensions. Another trade-off with increasing the NA is a decrease in the field of view. This is the same phenomenon experienced when going to a higher magnification with a zoom lens. At higher magnifications, the breadth of the view is narrowed. Field of view becomes a production limit with steppers. A narrower field requires more time to complete exposure of the entire wafer.

Other Exposure Issues

In [Chap. 8](#), the issue of image resolution and exposure wavelength was explored. In general, the way to expose smaller images is to use a smaller-wavelength exposure source. However, this leads to a smaller depth of field. In the sub-0.5- μm range, exposure processes extend from I-line to EUV. The DoF problem requires other refinements, including a variable NA lens, annular-ring illumination, off-axis illumination, and phase-shift masks. In addition, there are other optical effects that come into play as the image size gets smaller and the pattern density increases. These issues and solutions are examined next.

All of these technical problems are complicated by the increasing amount of information that must be printed on advanced circuits. With feature size shrinkage, chip size increases, stacking of components, and better designs, more information (per square centimeter) is required of the masking process.¹⁰ This trend pressures alignment makers, especially steppers, to develop ever-more sophisticated lens systems, exposure sources, resists, and other image enhancement techniques such as light source improvements, phase-shift masks (PSMs), and so forth (see [Chap. 10](#)). These techniques effectively lower the constant k in the resolution limit formula above.

Variable Numerical Aperture Lenses

The NA factor is an overall measurement of a lens's ability to gather light. The trade-off is depth of field (DoF). A lens with a better ability to gather light (higher NA) suffers from a decreased depth of field (see [Chap. 8](#)). Unfortunately, advanced circuits have many layers and high and low points (*surface topography*) that require a large DoF. A stepper with one lens is limited to exposing wafers within its associated depth of field, which may or may not include the levels on the wafer surface. Newer steppers come with a variable NA lens, allowing their use over a wider range of DoF requirements.^u

Immersion Exposure System

Light refraction is another factor in maintaining image fidelity to the mask or reticle dimensions. Diffraction is the bending of light around a mask edge or projection lens. The bending is a result of the wavelength of the light and index of refraction. It can be modified by flowing purified water between the lens and the wafer surface (see [Fig. 10.11](#)). The purified water, with a refractive index of 1.44, effectively increases the NA of the system and allows a smaller printed image size. Nikon claims that using an immersion system with an ArF¹ 193 nm excimer laser source, a 65-nm air image can be reduced to 40 to 45 nm¹² (see [Fig. 10.12](#)).

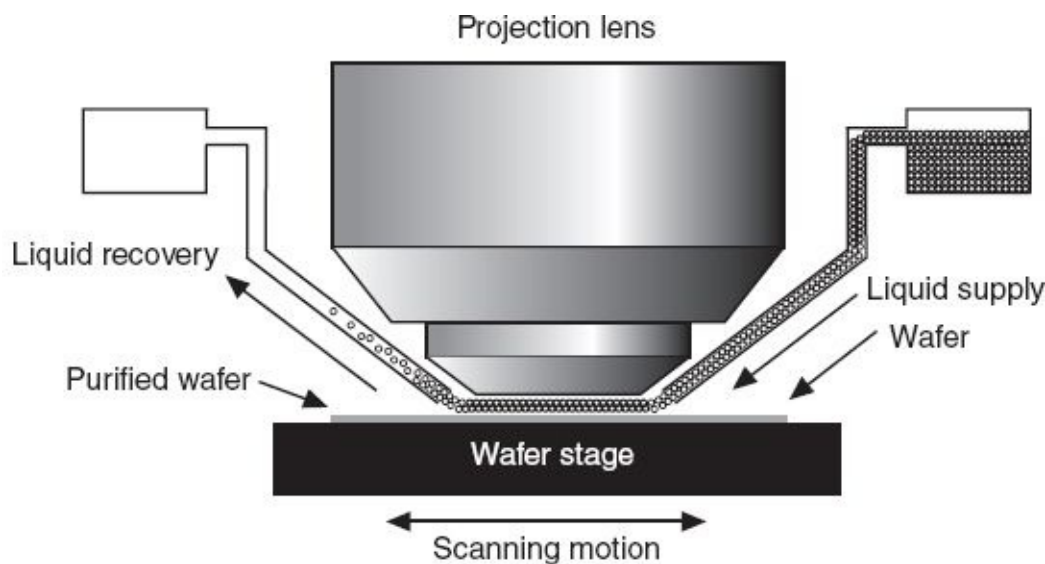


FIGURE 10.11 Immersion lens lithography tool.

	DRAM level			
	4 kb	64 kb	256 Mb	64 Mb → 1 GB
Aligner type	Contact	Projection	Stepper	Stepper-scanner
Exposure	Hg	Hg	g, I line	KrF, Arf
Image	Mask	Mask	Reticle	Reticle

FIGURE 10.12 Patterning techniques versus DRAM level.

Amplified Resist

A general resist resolution problem is the wavefront that arrives at the resist layer. Simple cross-sectional drawings show the arriving rays as uniform arrows. The actual radiation in the wavefront has a mixture of directions and energies that vary across surface and in the vertical direction. This is called the *aerial image*.¹³ Resolving one-half and one-third μm images with optical lithography requires management and manipulation of the aerial image. Control methods fall into three general areas: optical resolution, resist resolution limits, and surface problems. A fourth problem area is etch definition. Resist resolution is intimately mated to the exposure source and exposure system used. Basic resist formulation has improved with manufacturing control, matching the chemistry to the wavelength.

CA resists are considered the basic platform (with future improvements) to carry lithography to the 90-nm node and beyond. These resists, like older resists, are based on the photosensitive polymers, photoacid generators (PAGs), dissolution inhibitors, etch barriers, and an acid labile, base soluble group.¹⁴ New resists must operate in the tougher etch environments of plasma etch. They not only must image the smallest features, usually gates, but also deal with imaging dense patterns and small metallization contacts. Additionally, etch *line edge roughness* (LER) becomes a factor as line-size dimensions approach the size of the molecules in the resist.

Contrast Effects

Good resolution becomes difficult in regions of the mask where an opaque line is surrounded by a large clear area. The large amount of radiation coming around the opaque line tends to shrink the dimension of the line in the resist layer ([Fig. 10.13](#)) as the exposing rays diffract around the edge of the pattern. This problem is called a *proximity effect*.

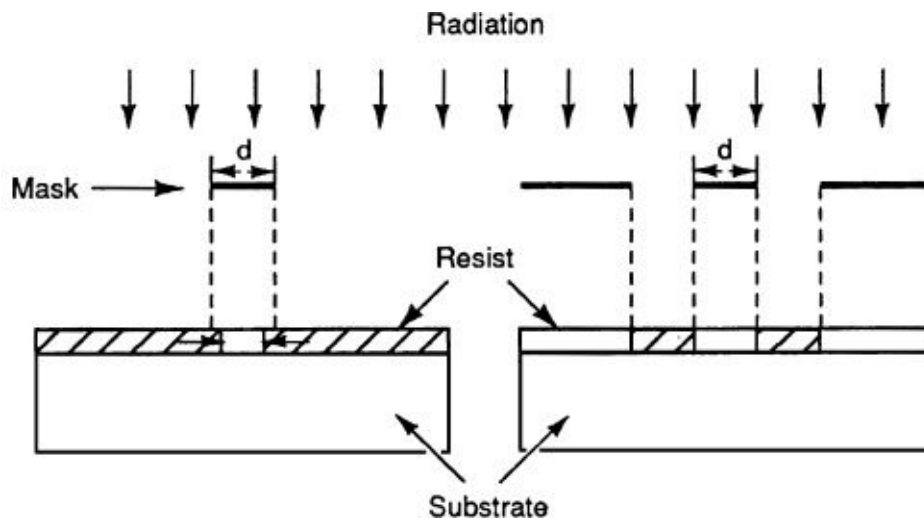


FIGURE 10.13 Proximity effects.

Another contrast effect is called *subject contrast* ([Fig. 10.14](#)). This situation comes about when some exposure radiation penetrates the opaque region of a mask or reflects off the wafer surface into the resist. The result is a partially exposed region that leaves a distorted image after the development step. This is more of a problem with negative resist than with positive resist. Image changes also occur from diffraction effects ([Fig. 10.15](#)) and light scattering ([Fig. 10.16](#)).

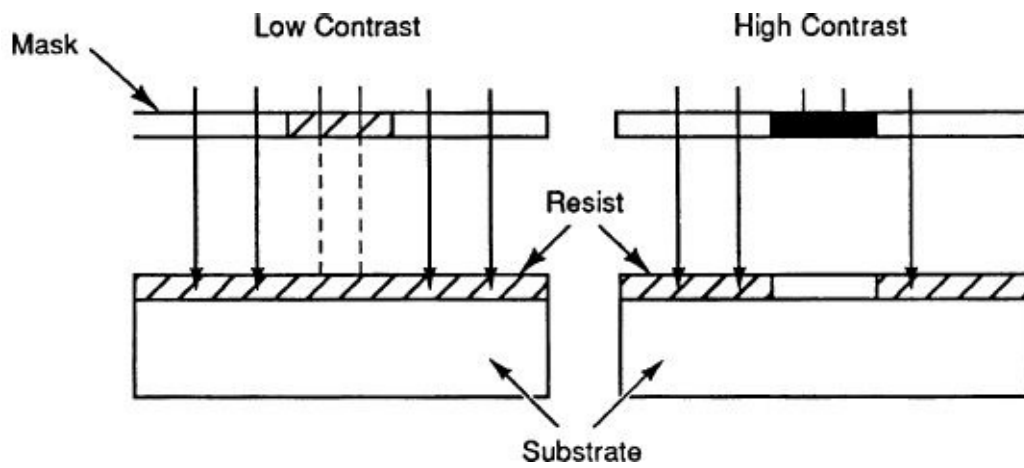


FIGURE 10.14 Subject contrast.

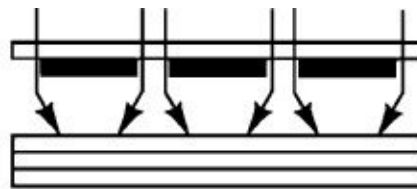


FIGURE 10.15 Diffracting reduction of image in resist.

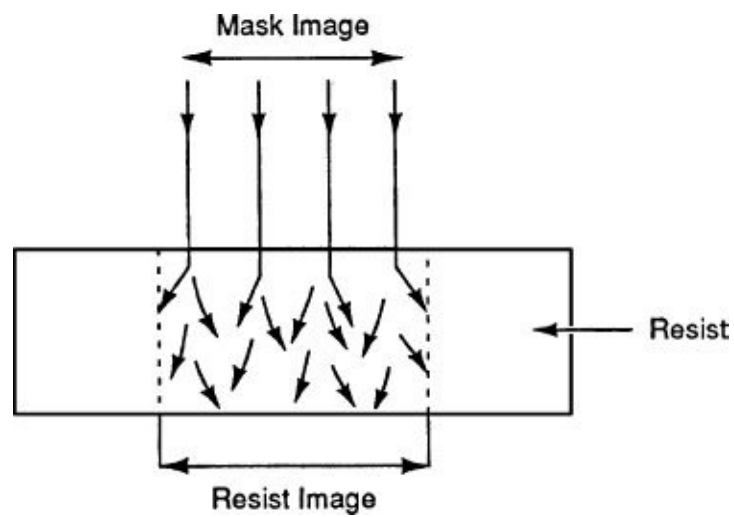


FIGURE 10.16 Light scattering in resist film.

Other Resolution Challenges and Solutions

In [Chap. 8](#), the issue of image resolution and exposure wavelength was explored. In general, the way to expose smaller images is to use a smaller-wavelength exposure source. However, this leads to a smaller depth of field. In the sub-0.5- μm range, exposure processes extend from I-line to deep UV. The DoF problem requires other refinements, including a variable NA lens, annular-ring illumination, off-axis illumination, and phase-shift masks. In addition, there are other optical effects that come into play as the image size gets smaller and the pattern density gets larger. These issues and solutions are examined next.

Off-Axis Illumination

Shifting the direction of the exposure beam from the perpendicular (off-axis) interrupts the interference pattern that causes standing waves in the resist.

Lens Issues and Reflection Systems

At the extremes of lithographic patterning shaping, the exposing beam through a lens system becomes an issue. The issue is absorption. An exposure system should deliver a specific wavelength (or bundle of controlled wavelengths) to the resist surface. The materials used for lenses can absorb radiation in the required ranges, and this becomes a serious problem below 193 nm. Calcium fluoride (CaF_2) is one material that is transparent in this radiation range, and it is expected to be used at the 157-nm node.¹⁵

Phase-Shift Masks

For conventional optical patterning, several techniques are used to improve image fidelity from mask to wafer. A diffraction problem occurs as two-mask patterns get closer together. At some point, the normal diffraction of the exposure rays start touching, leaving the patterns unresolved in the resist. The blending of the two diffraction patterns into one is because all the rays are in the same phase. *Phase* is a wave term that relates to the relative positions of the peaks and valleys of a wave ([Fig. 10.17a](#)). The waves in (a) are in phase, those in (b) are out of phase. One way to prevent the diffraction patterns from wiping out two adjacent mask patterns is to cover one of the openings ([Fig. 10.17b](#)) with a transparent layer that shifts one of the sets of exposing rays *out of phase*, which in turn nulls the blending. This approach to an *alternating phase-shift mask* (also called *alternating aperture phase-shift mask*—AAPSM) requires the deposition of a layer of silicon dioxide on the mask or reticle and a photomasking process to remove the oxide layer from alternate patterns. Covering every other clear opening works well for repeated array patterns such as those found in memory products.

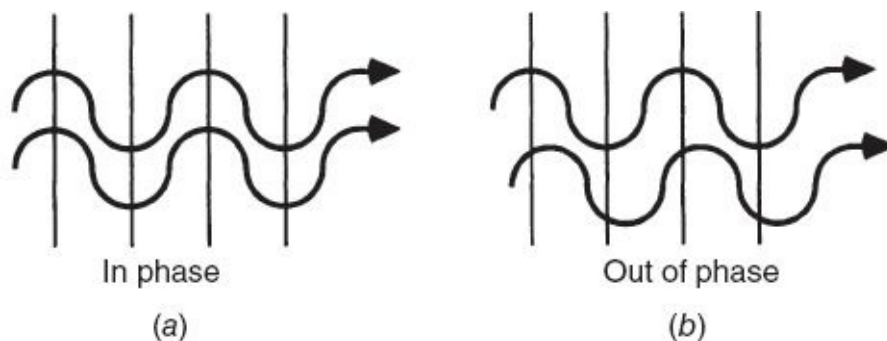


FIGURE 10.17 Wave phases.

Another solution is the addition of phase-shifting layers to the edges of the mask

or reticle patterns. This process also requires oxide deposition and full masking process. There are several variations for this approach. They are subresolution (or outrigger) and rim phase-shift masks¹⁶ (see [Fig. 10.18](#)).

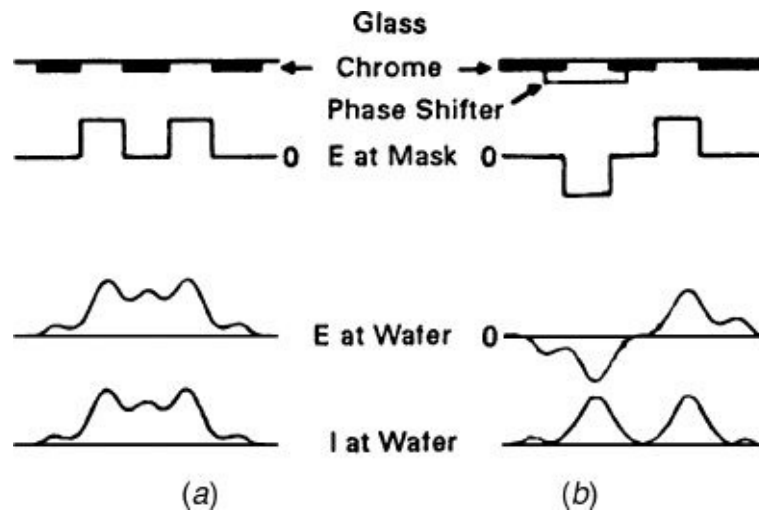


FIGURE 10.18 Light intensity patterns: (a) without phase shifting and (b) with phase shifting. (Source: VLSI Fabrication Principles, by *Ghandhi*.)

Optical Proximity Corrected or Optical Process Correction

It is known that, the sub-0.5- μm range can result in a distorted pattern in the resist. optical proximity corrected or optical process correction (OPC) masks attempt to reverse that situation by having a distorted image on the mask or reticle that is designed to produce a more faithful image in the resist. Situations that are particularly vulnerable are dense patterns of adjacent open or opaque regions and the shortening and/or rounding of pattern corners such as small contact holes. A computer is used to analyze exposure process conditions (contrast effects) and design the on-mask pattern shapes.¹⁷ An example of the correction of an end-rounded pattern is shown in [Fig. 10.19](#). Another technique is *double masking*. The first mask, a phase mask, creates part of the pattern in the resist. A follow-on mask, a trim mask, is somewhat oversized and completes the desired pattern.¹⁸

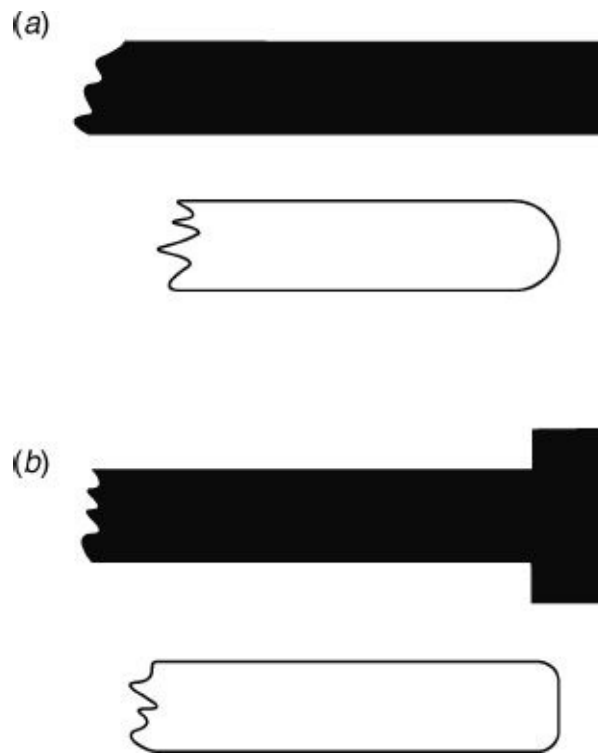


FIGURE 10.19 (a) Conventional image formation and (b) image enhancement with use of mask image “hammerhead.”

Annular-Ring Illumination

Annular-ring illumination is a technique that was first introduced in Perkin Elmer scanning projection aligners. One of the guns in the resolution arsenal is a more uniform exposing light. Unfortunately, conventional optical exposure sources produce a light spot that is too nonuniform for small image exposure. However, within the spot, there are areas (rings) of more uniform energy. An annular-ring illuminator blocks off all but a *ring* portion of the spot, directing a more uniform wave of exposing radiation to the wafer.

Pellicles

The development of projection exposure systems (projection aligners and steppers) brought with them an increased mask and reticle lifetime. With the increased lifetime came the incentive to make higher-quality masks and reticles. In a production line where masks are used for a long time, wafer-sort yield loss comes from dirt and scratches picked up during handling and use. One source of damage comes from mask and reticle cleaning steps. This situation is a “Catch-22.” The masks and reticles become dirty during the process and require cleaning. The cleaning procedure then itself becomes a source of contamination, scratches, and breakage.

A solution to these problems is a pellicle ([Fig. 10.20](#)). A *pellicle* is a thin layer of an optically neutral polymer stretched onto a frame. The frame is designed to fit onto the mask or reticle. The pellicle is fitted to the mask or reticle after the mask is made and cleaned. Once in place, the pellicle membrane is the surface collecting any dirt or dust in the environment. The height of the membrane above the mask surface holds the dirt particles out of the focal plane of the mask. In effect, the particles are transparent to the exposing rays.

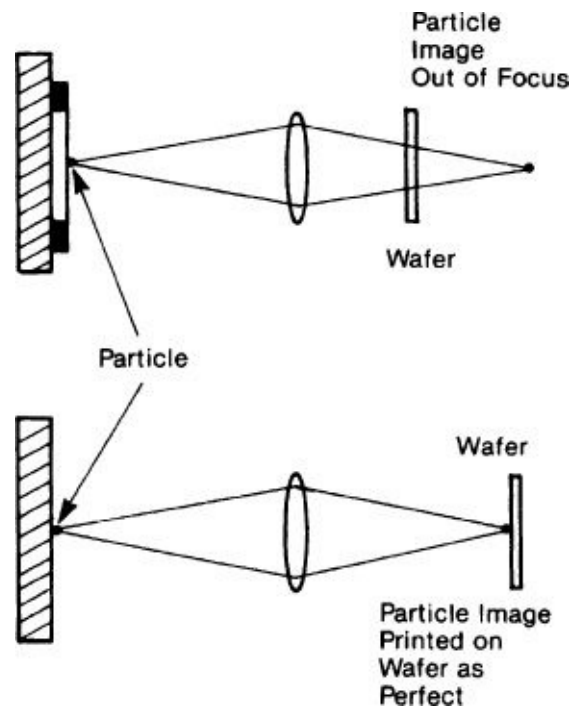


FIGURE 10.20 Pellicle.

Another benefit of a pellicle is the elimination of scratches from the mask surface, since the surface is covered. A third benefit is that a pelliclized, mask or

reticle does not need in-process cleaning. For some applications, the pellicle membrane is given an antireflective coating, which assists in the imaging of small geometries, especially on reflective wafer surfaces. These benefits can translate into wafer-sort yield increases of 5 to 30 percent.¹⁹

Pellicle membranes are made from nitrocellulose (NC) or cellulose acetate (AC). NC films are used in exposure systems with broadband exposure sources (340 to 460 nm)²⁰ while AC films are used in mid-ultraviolet applications. In the 248-to 193-nm range Teflon AF[®] or Cytop[®] are used.²¹ The membranes are thin (0.80 to 2.5 μm) and must show a high-transmission rate for the exposure wavelengths. A typical pellicle will exhibit over a 99 percent transmission rate for the peaks of the exposing wavelengths. Pellicle effectiveness requires stringent thickness control, on the order of $\pm 800 \text{ \AA}$, and control of particles to less than 25 μm in diameter.

A pellicle membrane is made by a spin-casting technique. The pellicle material is dissolved in a solvent and spun onto a rigid substrate, such as a glass plate. This is the same technique used to spin photoresist onto wafers. Thickness of the membrane is controlled by the viscosity of the solution and the spin speed of the spin coater. The membrane is removed from the substrate and fixed onto the frame. Frame shapes are determined by the size and shape of the mask or reticle. Cleanliness control requires Class 10 or better cleanrooms and antistatic packaging.

Surface Problems

The resolution of small images is affected by several conditions on the wafer surface. Reflections off the surface layers, increasing variation of the topography, and the etching of multilayer stacks all require special process steps or “tweaking” of the process.

Resist Light Scattering

In addition to light radiation reflecting off the wafer surface, the radiation tends to diffuse into the resist-causing poor image definition. The amount of diffusion is in proportion to the resist thickness. Some additives put in the photoresist to increase radiation absorption also increase the amount of radiation diffusion, thus reducing image resolution.

Subsurface Reflectivity

The high-intensity exposing radiation ideally is directed at a 90° angle to the wafer surface. When this ideal situation exists, exposing waves reflect directly up and down in the resist, leaving a well-defined exposed image ([Fig. 10.21](#)). In reality, some of the exposing waves are traveling at angles other than 90° and expose unwanted portions of the resist.

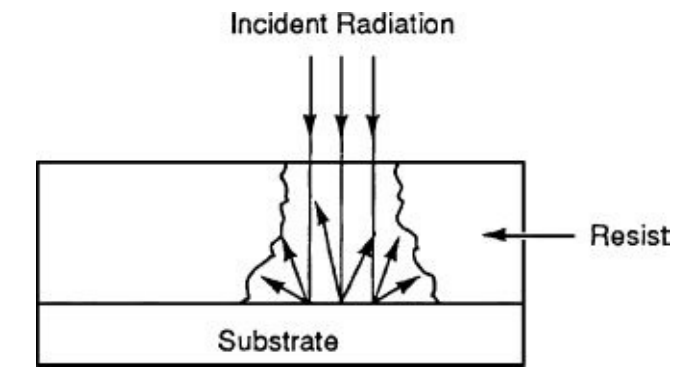


FIGURE 10.21 Subsurface reflectivity.

This subsurface reflectivity varies with the surface layer material and the surface smoothness. Metal layers, especially aluminum and aluminum alloys, have higher reflectivity properties. A goal of the deposition processes is a consistent and smooth surface to control this form of reflection.

Reflection problems are intensified on wafers with many steps, also called a varied topography. The sidewalls of the steps reflect radiation at angles into the resist, causing poor image resolution. A particular problem is light interference at the step that causes a “notching” of the pattern as it crosses the step ([Fig. 10.22](#)).

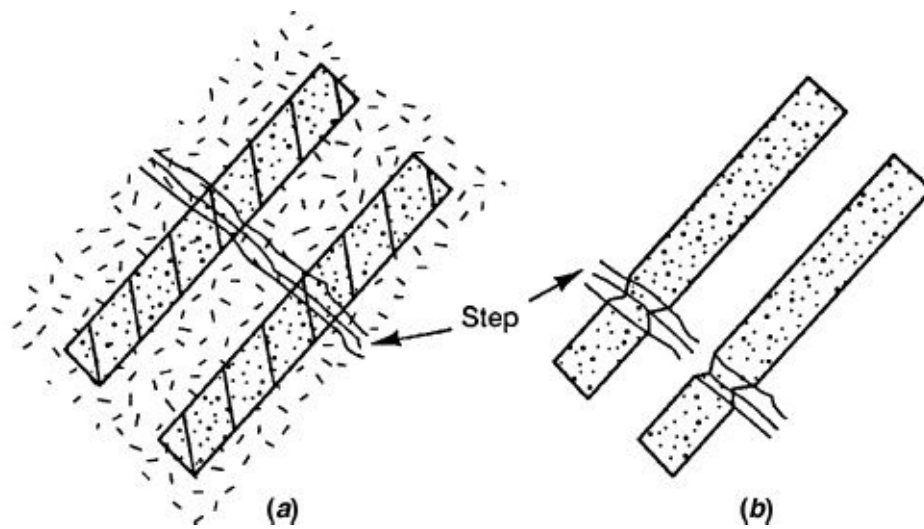


FIGURE 10.22 Metal line notching over a step: (a) before etching; (b) after etching.

Antireflective Coatings

Antireflective coatings (ARCs) spun onto the wafer surface before the resist ([Fig. 10.23](#)) can aid the patterning of small images. The ARC layer brings several advantages to the masking process. First is a planarizing of the surface, which in turn makes for a more planarized resist layer. Second, an ARC cuts down on light scattering from the surface into the resist, which helps in the definition of small images. An ARC can also minimize standing wave effects and improve the image contrast. The latter benefit comes from increased exposure latitude with a proper ARC.

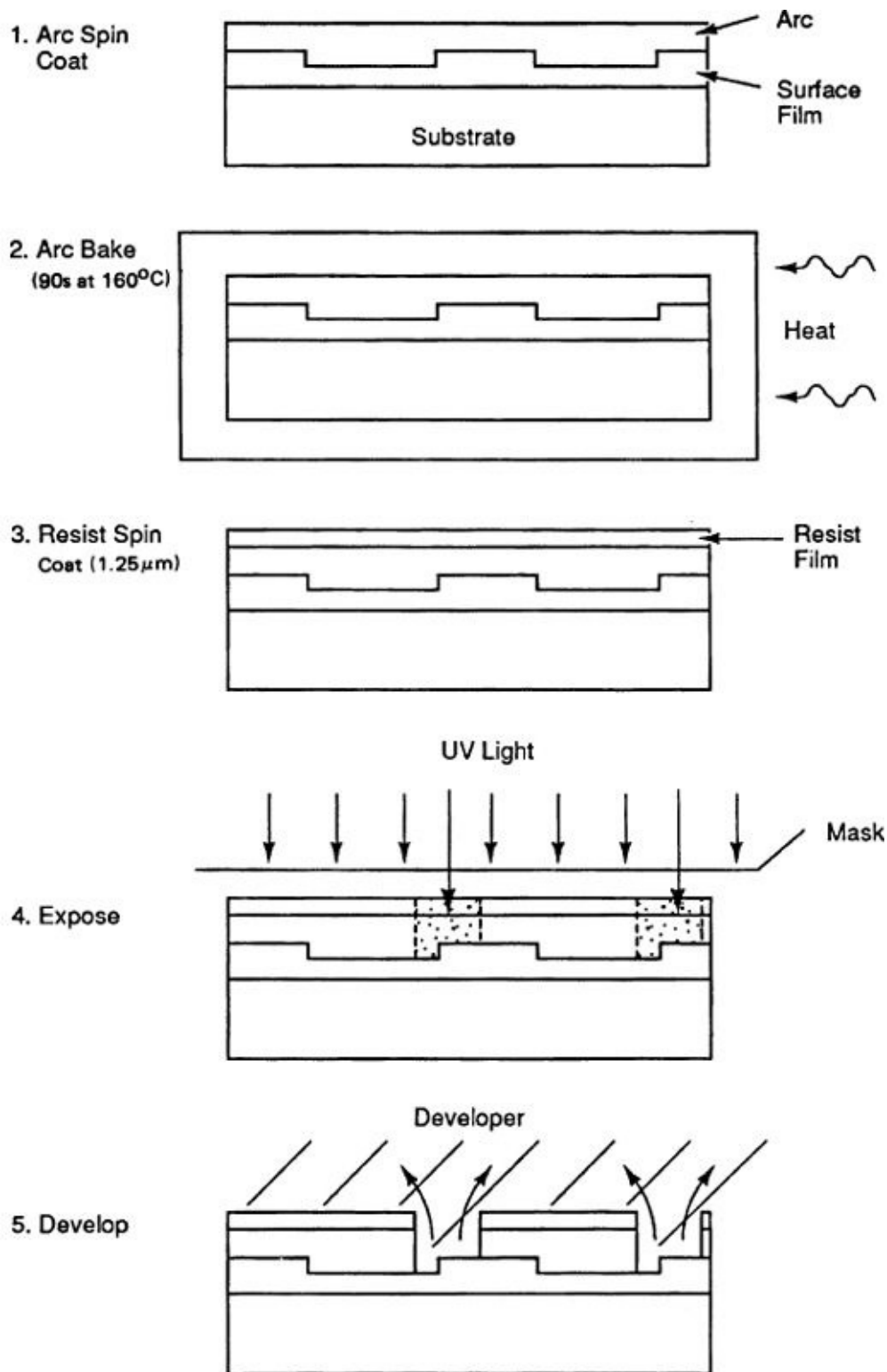


FIGURE 10.23 Antireflective process sequence.

An ARC is spun onto the wafer and baked. After the resist is spun on top of the ARC, the wafer is aligned and exposed. The pattern is developed in both the resist and the ARC. During the etch, the ARC acts as an etch barrier. To be effective, an ARC material must transmit light in the same range as the resist. It must also have

good adhesion properties with the wafer surface and the resist. Two other requirements are that the ARC must have a refractive index that matches the resist, and the ARC must develop and be stripped in the same chemicals as the resist.

There are several penalties associated with the use of an ARC. One is an additional layer requiring a separate spin and bake. The resolution gains offered by an ARC can be offset with poor thickness control and/or an ill-controlled developing step. The time of exposure can increase 30 to 50 percent increasing the wafer throughput time. ARC layers may also be used as the intermediate layer in a trilayer resist process or used on the top of the photoresist (top antireflective coating, or TAR).

Standing Waves

In the “Subsurface Reflectivity” section, it was mentioned that the ideal exposure situation is when the radiation waves are directed to the wafer surface at 90° . This is true when only reflection problems are under consideration. However, 90° reflection causes another problem in positive photoresists—the creation of standing waves. As the radiation wave reflects off the surface and travels back up through the resist, it interferes constructively and destructively with the incoming wave, creating regions of varying energy ([Fig. 10.24](#)). The result, after development, is a rippled sidewall and a loss of resolution. A number of solutions are used to moderate standing waves, including dyes in the resist and separate antireflective coatings directly on the wafer surface. Most positive resist processes include a post-exposure bake (PEB) step before development of the resist layer. The bake reduces the influence of standing waves on pattern sidewall definition.

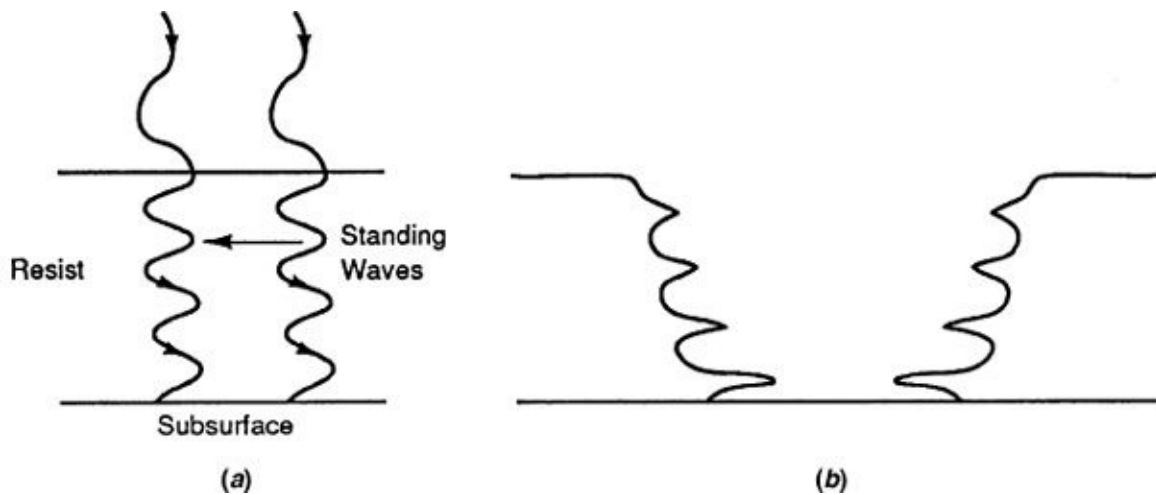


FIGURE 10.24 Standing-wave effect: (a) during exposure and (b) after develop.

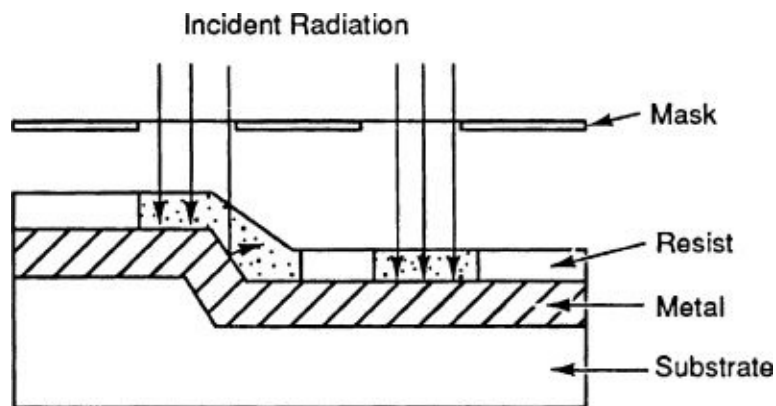


FIGURE 10.25 Light reflection at steps.

Planarization

The advancement of circuits to VLSI and ULSI levels has resulted in both increased surface density and more surface layers. As the various layers are etched into patterns, the surface becomes stepped into high and low plateaus. The plateaus are delineated by steps each with different reflection qualities ([Fig. 10.24](#)). The resolution of submicron images on this type of surface is not possible in a simple one-resist layer process. One problem is depth-of-field issues. A limited-DoF lens cannot resolve images on both the high and low plateaus. Another is light reflection at the steps causing the notching of metal lines ([Fig. 10.22](#)). A number of *planarization* techniques are used to offset the effects of a varied wafer topography. Techniques include multilayer resist processing, planarization layers, reflow techniques, and physically flattening the surface by chemical-mechanical polishing (CMP). Current focus is on the global flattening of the surface (CMP) to allow patterning with only one resist layer. However, the multilayer resist processes described are still useful during the transition.

Photoresist Process Advances

The ten-step lithography process detailed in [Chaps. 8](#) and [9](#) is based on a single photoresist layer, and it assumes that the layer can resolve the required images without pinholes or failing during the etch process. The nanometer era has forced the development of variations to the single-resist layer process. These processes are usually used in conjunction with the advances in exposure control described previously.

Multilayer Resist or Surface Imaging

There are a number of multilayer resist processes. The choice of a process depends on the size of the resist opening and the severity of the surface topography. While multilayer resist processes add the penalty of more process steps in some situations, they are the only reliable means of creating the desired images. A multilayer resist process features a thicker bottom layer that serves to fill in the valleys and planarize the surface. The image is first formed in a layer of photoresist on the top of the planarizing layer(s). This *surface imaging* allows small dimension imaging because the surface is flat, that is, away from reflecting steps and in the same plane to avoid depth of focus (DOF) problems.

A dual multilayer resist process uses two layers of photoresist, each with a different polarity. The process is suited to resolve small geometries on wafers with a varied topography. First, a relatively thick layer of resist is applied and baked to the thermal flow point ([Fig. 10.26](#)). A typical thickness is three to four times the highest step height on the wafer. The goal is a planar top resist surface. A typical multilayer process will use a positive-acting polymethylmethacrylate (PMMA) resist that is sensitive to deep ultraviolet radiation.

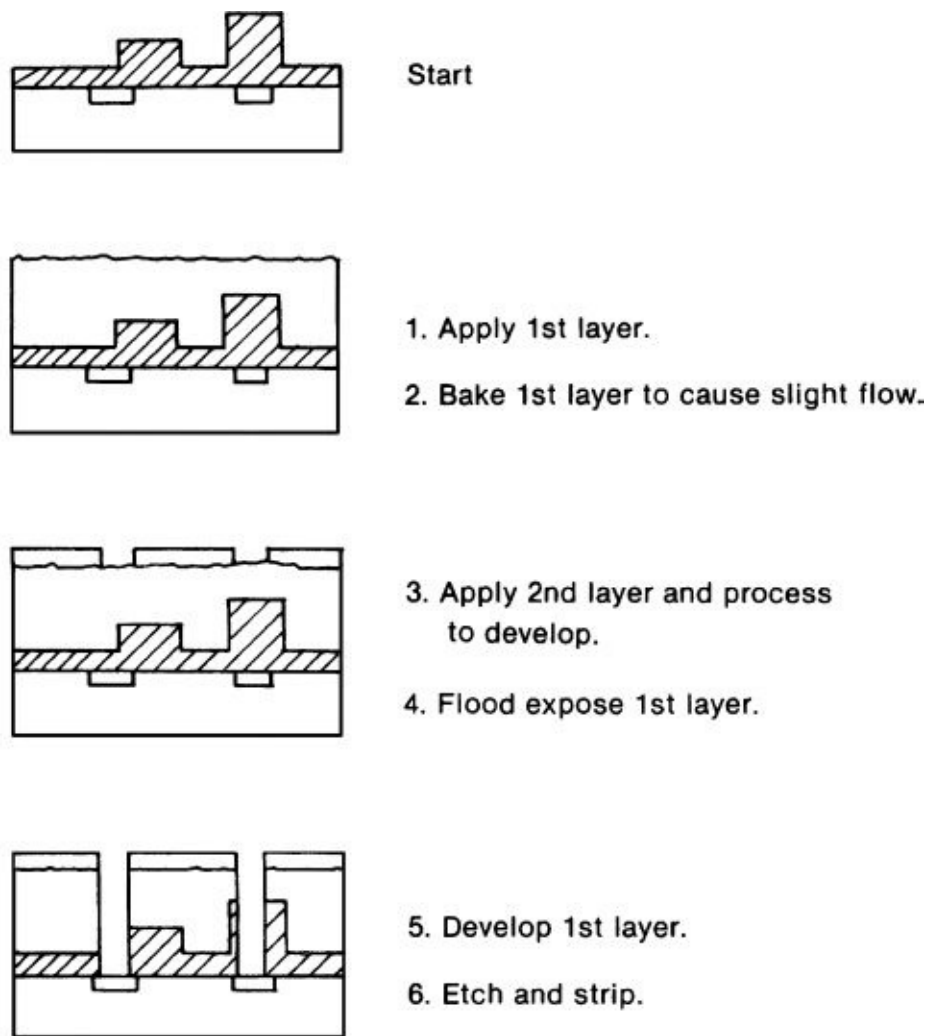


FIGURE 10.26 Dual-layer photoresist processing.

Next, a thin layer of positive resist, sensitive only to ultraviolet radiation, is spun on top of the first layer and processed through the development step. The thin top layer allows the resolution of the pattern without the adverse effects encountered with thick resist layers or reflections from steps in the surface. Since the top layer conforms to the shape of the bottom layer, it is referred to as a *conformal layer* or *portable conformal layer*. This top layer of resist acts as a radiation block, leaving the bottom layer unpatterned. Next, the wafer is given a blanket or flood (no mask) deep ultraviolet exposure, which exposes the underlying positive resist through the holes in the top layer, thus extending the pattern down to the wafer surface. A development step completes the hole resolution and the wafer is ready for etch.

Considerations in the choice of photoresists are compatibility of the two resists through the process, reflection problems from the subsurface, standing waves, and sensitivity problems with PMMA resists.²² In addition, the two resists must have compatible bake processes and independent developing chemistries.

Variations on the basic dual-level resist process include dyes in the PMMA and

the use of antireflection layers under the first resist layer. Many variations on the basic dual-level process are practiced. One use of a dual-layer resist process is as a lift-off technique. By adjusting the development of the bottom layer, an overhang can be created that assists in the clean definition of the metal line on the surface ([Fig. 10.27](#)).

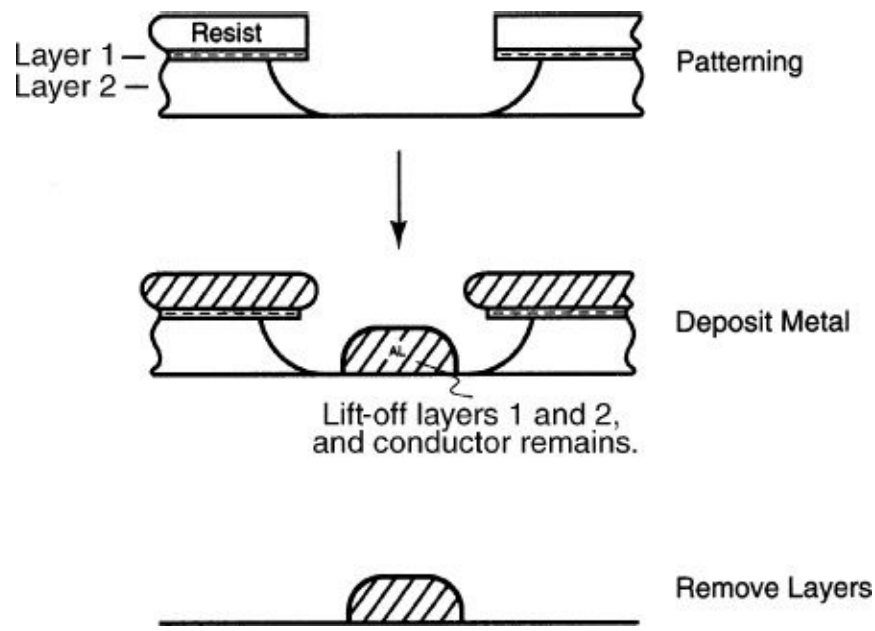


FIGURE 10.27 Dual-layer resist lift-off process.

A trilevel resist process ([Fig. 10.28](#)) incorporates a “hard” layer between the two resist layers. The hard layer may be a deposited layer of silicon dioxide or other developer-resistant material. As in the dual-layer process, the image is formed in the top photoresist layer. Then, the image is transferred into the hard layer by a conventional etching step. The finishing step is the formation of the pattern in the bottom layer, using the hard layer as an etch mask. The use of a hard intermediate layer makes possible the use of nonphotoresist bottom layers such as a polyimide layer.

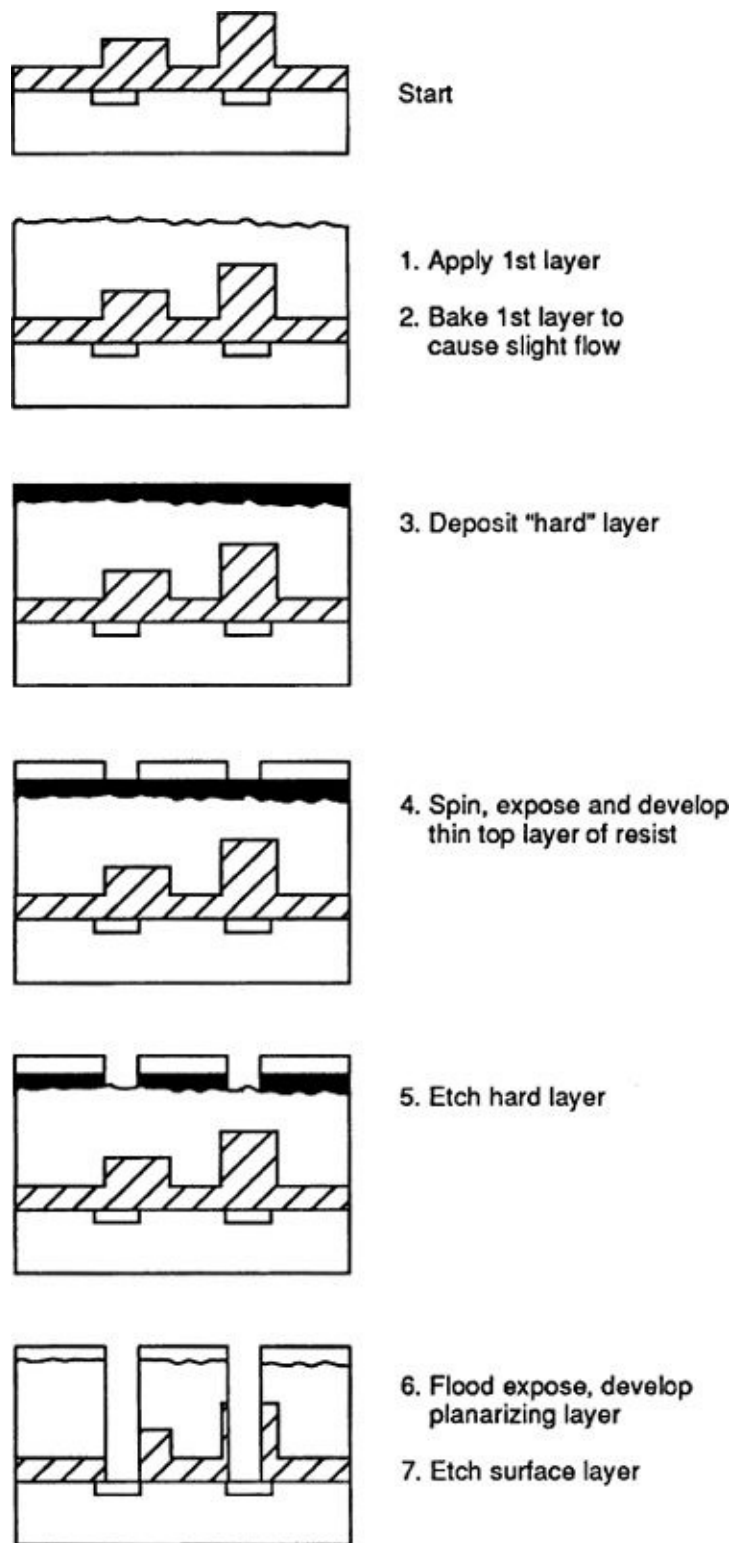


FIGURE 10.28 Trilevel resist process.

Silylation or DESIRE Process

A novel approach to surface imaging is the diffusion-enhanced silylating resist (DESIRE) process²³ ([Fig. 10.29](#)). Like other multilayer resist processes, the concept is to planarize the wafer and form the image in a surface layer. The DESIRE process uses one layer exposed by a standard UV exposure. In this process, the exposure is confined to the top layer of the planarizing layer. Next, the wafer is placed in a chamber (see vapor prime baking)²¹ for exposure to HMDS for a *silylation process*. During this step, silicon becomes incorporated into the exposed areas. The silicon-rich areas become a hard mask, allowing the dry development and removal of the underlying material with an anisotropic RIE etch ([Chap. 9](#)). During the etch step, the silylated areas are converted to silicon dioxide (SiO_2), forming a more resistant etch mask. Techniques relying on defining the pattern in the topmost layer are called *top surface imaged* (TSI).

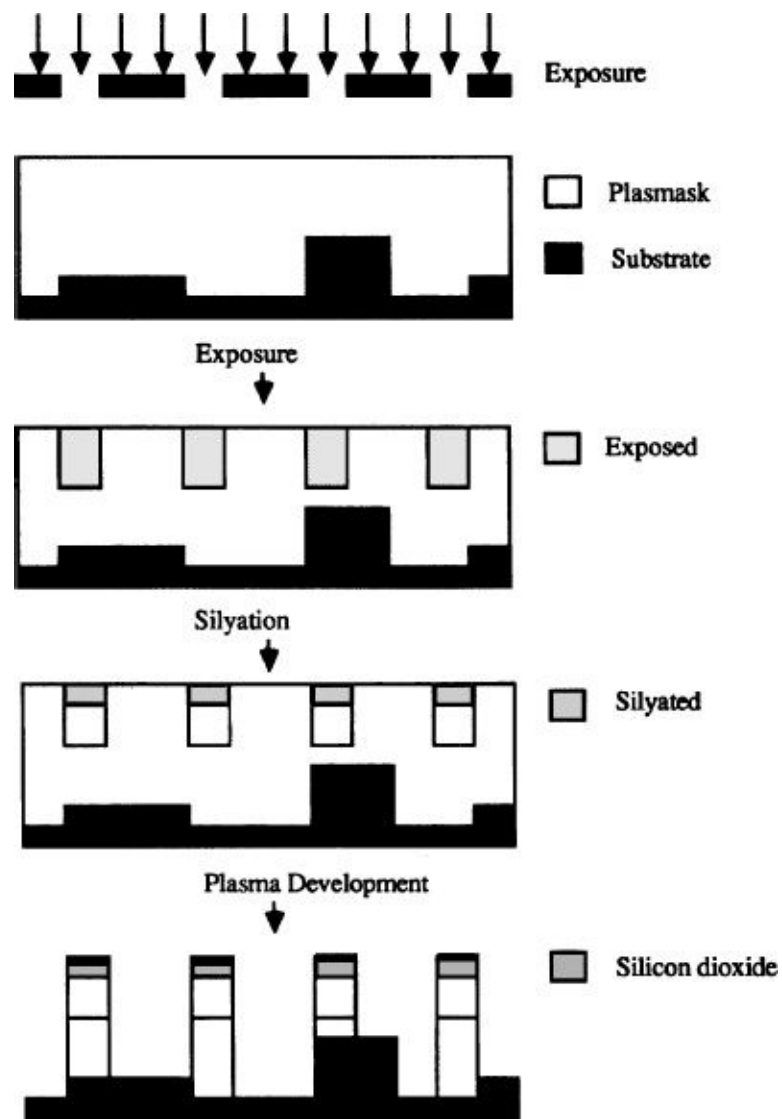


FIGURE 10.29 DESIRE process. (Source: Solid State Technology, June 1987.)

Polyimide Planarization Layers

Polyimides have been used for years in printed circuit board manufacturing. For semiconductor use, the polyimides offer the dielectric strength of deposited silicon dioxide films and the process advantage of application to the wafer with the same spinning equipment used for photoresist.²⁴

Once applied to the wafer, the polyimides flow over the surface, making it more planar. After application and flow, the polyimide can be covered with a hard layer and patterned with chemicals much like a photoresist. A popular use of polyimide layers is as an interdielectric layer between two layers of conducting metal. The planarizing effect of the polyimide makes the definition of the second metal layer easier.

Etchback Planarization

Etchback is used for local planarization ([Fig. 10.30](#)). After metal lines are defined, a thick oxide layer is deposited and a photoresist layer spun on top of the oxide. Etch takes place in a plasma etcher. First, the thinner resist etches away and starts etching the oxide. Later, the thicker resist etches away and some oxide is removed. The net result is a local flattening of the surface.

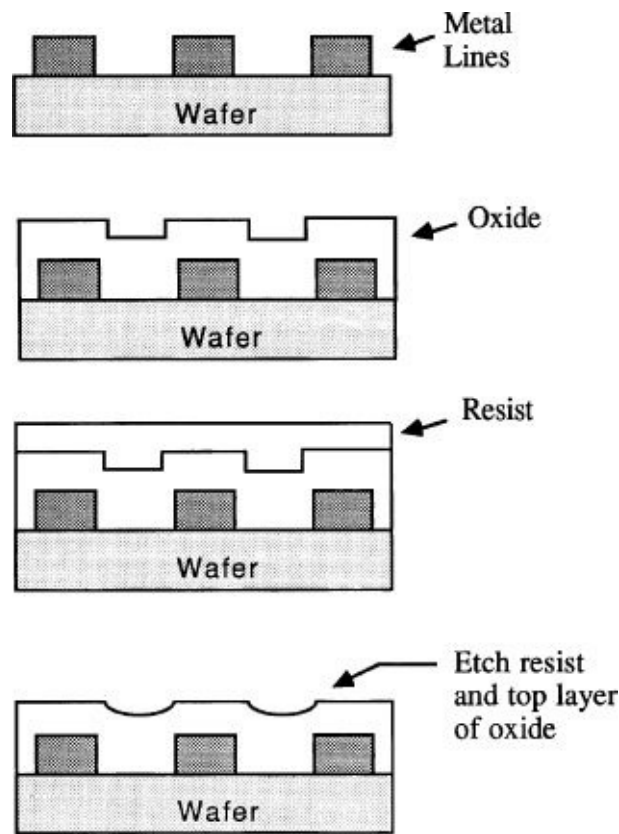


FIGURE 10.30 Etchback planarization.

Dual-Damascene Process

Increased component density has forced the use of multimetallization layers. Their use required the need for connecting conductors called *studs* or *plugs*. Tungsten is the preferred metal, but there are complications in etching tungsten. Also, copper has become the preferred metallization system, replacing aluminum. However, copper technology introduces a whole host of process issues. One is the replacement of the traditional image and etch processes with a process called *dual-damascene*. It is an inlaid process similar in principle to the inlaid metal processes of antiquity. In that process, grooves are made in the surface of a bowl or other object. A metal is

applied to the entire surface, also filling the groove. After the excess metal on the surface is removed, some remains in the grooves, leaving a decorative pattern. In the semiconductor application, the grooves are created with lithography techniques and the copper is deposited by electroplating. The metal deposition process overfills the surface. A *chemical mechanical polishing* (CMP) step is used to remove the overspill, leaving the metal isolated in its trench ([Fig. 10.31](#)). [Chapter 13](#) has a more detailed discussion on this new and important patterning technique.

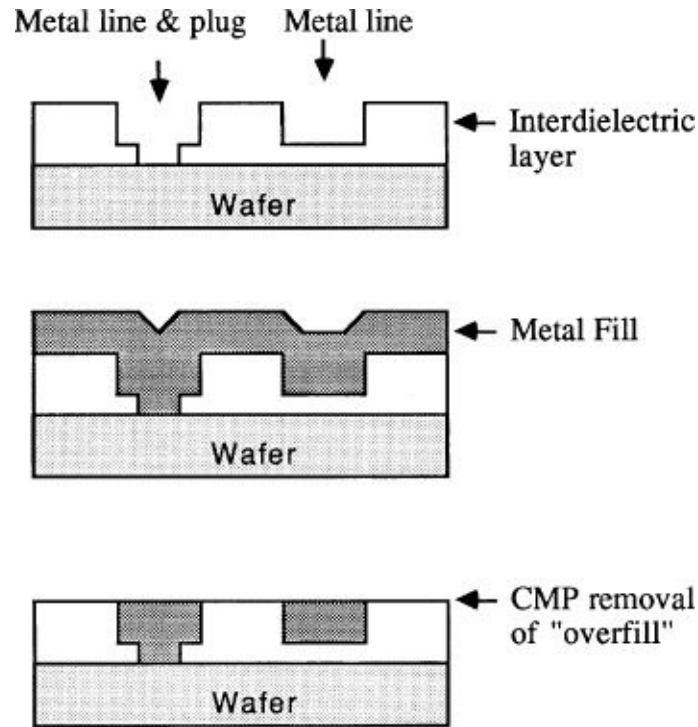


FIGURE 10.31 Dual-damascene (inlaid) process.

Chemical Mechanical Polishing

Each of the previously described planarization techniques tends to smooth out the surface but does not produce a totally flat surface (*global planarization*). Small-dimension images with most planarization techniques are still difficult to define because of light scattering. They also can cause metal-step coverage problems at the steps. The chemical-mechanical polishing process used to flatten wafers in the wafer preparation stage ([Chap. 3](#)) is also used for the global planarization of in-process wafer surfaces.

CMP is favored because of its ability to²⁵:

- Achieve *global planarization* across the entire wafer

- Polish and remove all materials
- Work on multimaterial surfaces
- Make possible a high-definition damascene process and copper metallization
- Avoid the use of hazardous gases
- Be a low-cost process

The same basic process is used for both wafer polishing and planarization ([Fig 10.32](#)). However, the challenges are very different. In wafer polishing, several microns of silicon are removed. In metallization processes, CMP is required to remove amounts of material on the order of a micron or less. Also, these metal processes present a number of materials for removal in the same step. They have different removal rates and challenges for uniform planarization. These issues are addressed in [Chap. 13](#).

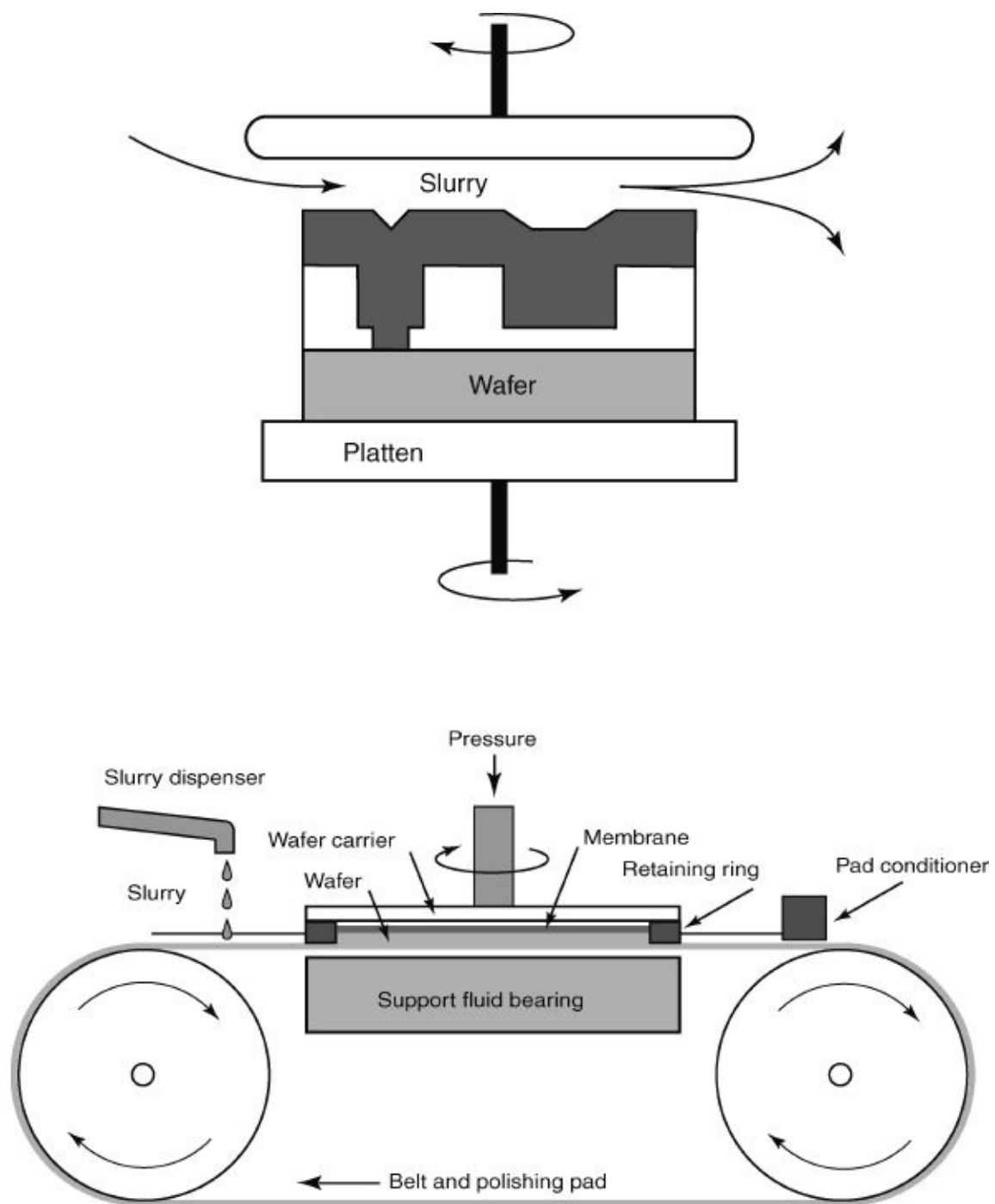


FIGURE 10.32 Chemical-mechanical polishing and planarization.

Basic CMP Processing Steps

A wafer is mounted upside down on a carrier, which, in turn, is mounted, wafer face down, onto a rotating platen. The platen surface is covered with a polishing pad. *Slurry*-carrying small abrasive particles are flown onto the platen. The particles attack and remove small pieces of the wafer surface, which are carried away by the movement of the slurry across the surface. The combined actions of the two (orbital) rotations and the abrasive slurry polish the wafer surface. High plateaus on the wafer are polished first and faster than the lower areas, thus achieving planarization. These are the mechanical polishing actions. However, mechanical polishing alone is unacceptable for semiconductor processing due to excessive mechanical damage to the surface. Damage is reduced and/or managed by selecting a slurry chemistry that dissolves or etches the surface materials. Chemical removal generally requires the corrosion of the surface, usually through an oxidation mechanism. An analogy would be the rusting of iron, which is chemical corrosion. A layer of rust forms when iron is exposed to oxygen. However, the layer of rust slows down the rusting process by blocking the iron surface from the oxygen in the water or air. This is where the chemical and the mechanical work together. The abrasive particles wipe away the corrosion layer, exposing a fresh surface for corrosion, and the process repeats itself continuously. Following CMP, there is a *post-CMP cleaning* required to maintain wafer cleanliness.

Linear CMP systems are also used. In this arrangement, the wafer carriers are rotating over a moving belt instead of a rotating platen. A chief advantage of linear systems is the increased speed of the slurry movement under the wafer.

Primary measures of performance after a CMP process are:

- The surface flatness
- Surface mechanical condition
- Surface chemistry
- Surface cleanliness
- Productivity
- Cost of ownership (see [Chap. 15](#))

CMP Polishing Pads

Polishing pads are made of cast polyurethane foam with fillers, polyurethane impregnated felts, or other materials with special properties. Important pad properties are porosity, compressibility, and hardness.^{[26](#)} Porosity, usually measured as the specific gravity of the material, governs the pad's ability to deliver slurry in its pores and remove material with the pore walls. Compressibility and hardness relate to the pad's ability to conform to the initial surface irregularities. Generally, the harder the pad, the more global the planarization. Softer pads tend to contact both the high and low spots, causing nonplanar polishing. Another approach is flexible polish heads that allow more conformity to the initial wafer surface.^{[27](#)}

Slurry

Slurry chemistry is complex and critical as a result of its dual role. On the mechanical side, the slurry is carrying abrasives. Small pieces of silica (silicon dioxide) are used for oxide polishing. Alumina (Al_2O_3) is a standard for metals. Multimetal surfaces found in multimetal layer schemes ([Chap. 13](#)) are challenging the industry to identify more “universal” abrasives. Abrasive diameters are kept in the 10-to 300-nm size²⁸ to achieve polishing, as opposed to grinding, which uses larger-diameter abrasives but causes more surface damage.

On the chemical side, the etchant would be KOH or NH_4OH (basic solutions with low pH levels) for silicon or silicon dioxide. For metals such as copper, reactions usually start with an oxidation of the metal from the water in the slurry. A typical reaction is: $2\text{Cu} + \text{H}_2\text{O} \rightarrow \text{Cu}_2\text{O} + 2\text{H}^+ + 2\text{e}^-$

Following the oxidation, the basic materials chemically reduce the film, which is removed by the mechanical actions. Various additives are found in production slurries. They perform various roles. Reducing post-CMP surface residues is addressed by balancing the pH of the slurry to control electrical charges on the²⁹ abrasive particles. Silica-based slurries have high pH levels, while silica slurries have a pH below 7. Other additives are surfactants to establish desired flow characteristics and chelating agents. These latter agents interact with metal particles to reduce their redeposition on the wafer surface.

Other factors critical to a planar polish are the pH of the slurry (degree of acidity or alkalinity), the flow dynamics at the wafer-pad interface, and the etch selectivity of the slurry on different surface materials and underlying layers.

Polishing Rates

A primary production measure is the polishing rate; many factors influence the rate. Pad material parameters already described, the slurry types and size, and the chemicals and their properties used for corroding the surface are important factors. Other process factors include the pad pressure, the rotation rates, the flow rate of the slurry, the flow property (viscosity) of the slurry, and the temperature and humidity in the polishing chamber. Wafer diameter, pattern sizes, and surface materials are also factors. All of these factors must be balanced to achieve a productive polish rate (go faster) without creating an out-of-control process (go slower).

Planarity

Global planarity is the goal of CMP, but the advent of multimetal schemes challenges that goal. Copper is a particular issue and illustrates several of the basic problems. Copper deposition into the trenches of a damascene patterning system ([Fig. 10.33](#)) results in lower density in the center. During the CPM process, the center polishes faster leaving a dish shape. Also, copper deposition density differences can result in differential polishing across the pattern.

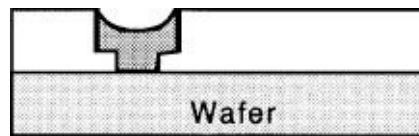


FIGURE 10.33 “Dishing” of copper in trench.

Tungsten plugs are also a CMP challenge ([Fig. 10.34](#)). During the initial CMP, the tungsten surface ends up recessed below the surrounding oxide. An oxide buff is required to planarize the surface.³⁰

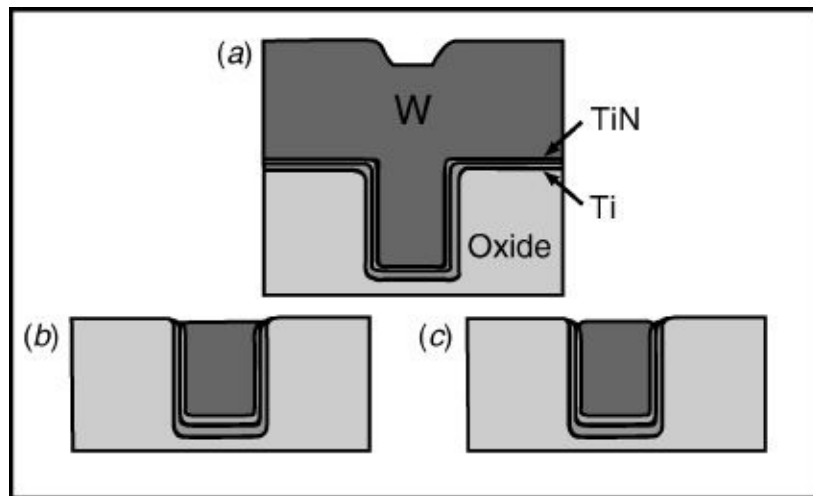


FIGURE 10.34 Tungsten plug formation: (a) deposit W, (b) CMP removal, and (c) buff oxide layer. (Reprinted from the April edition of *Solid State Technology*, Copyright 1998 by PennWell Publishing Company.)

In some copper metallization schemes, tantalum is used as a diffusion barrier in trenches to keep the copper from diffusing into the silicon. However, the tantalum polishes much more slowly than copper (in a copper-oriented slurry), leaving the copper surface exposed to more polishing time and more dishing.³¹

Pattern geometry variations result in differential removal rates. Larger areas tend to be removed faster, leaving dishing of low spots on the surface.

Also challenging is the presence of metals of different hardness, which polish at different rates, and interdielectric layers (IDLs) of polymer materials that are soft. Nevertheless despite all the challenges, wafer surfaces have to be flat down to the 150-nm range or lower.

Post-CMP Clean

The critical role of clean wafer surfaces has been stressed throughout this book. It is just as important after CMP, and there are some particular challenges. CMP is the only process that intentionally puts particulates on the surface, namely, the abrasive particles. These are usually removed with mechanical brush cleaners or high-pressure water jets. Chemical cleaning generally employs the same techniques used for other FEOL cleaning.

Carefully choosing slurry surfactants and adjusting the pH can create an electrical repulsion between the slurry particles and the wafer surface. This technique can reduce contamination, particularly of the types that become electrostatically attached to the surface.

Copper contamination is a particular concern. If it gets into the silicon, it changes and diminishes electrical operations of the circuit components. Copper residues should be reduced to the 4×10^{13} atoms/cm² range.^{[32](#)}

CMP Tools

Operating a successful CMP process requires a sophisticated integrated system more than the simple polishing unit ([Fig. 10.35](#)). Production-level tools include automatic wafer-handling robots, on-board metrology, and cleanliness detectors. Various end-point detection systems are used to signal when a particular layer material is gone or a specific removal depth is reached. Post-CMP cleaning units may be included in the main cabinet or mated to the primary CMP unit by robots in a cluster design. The goal is a “dry-in, dry-out” process.

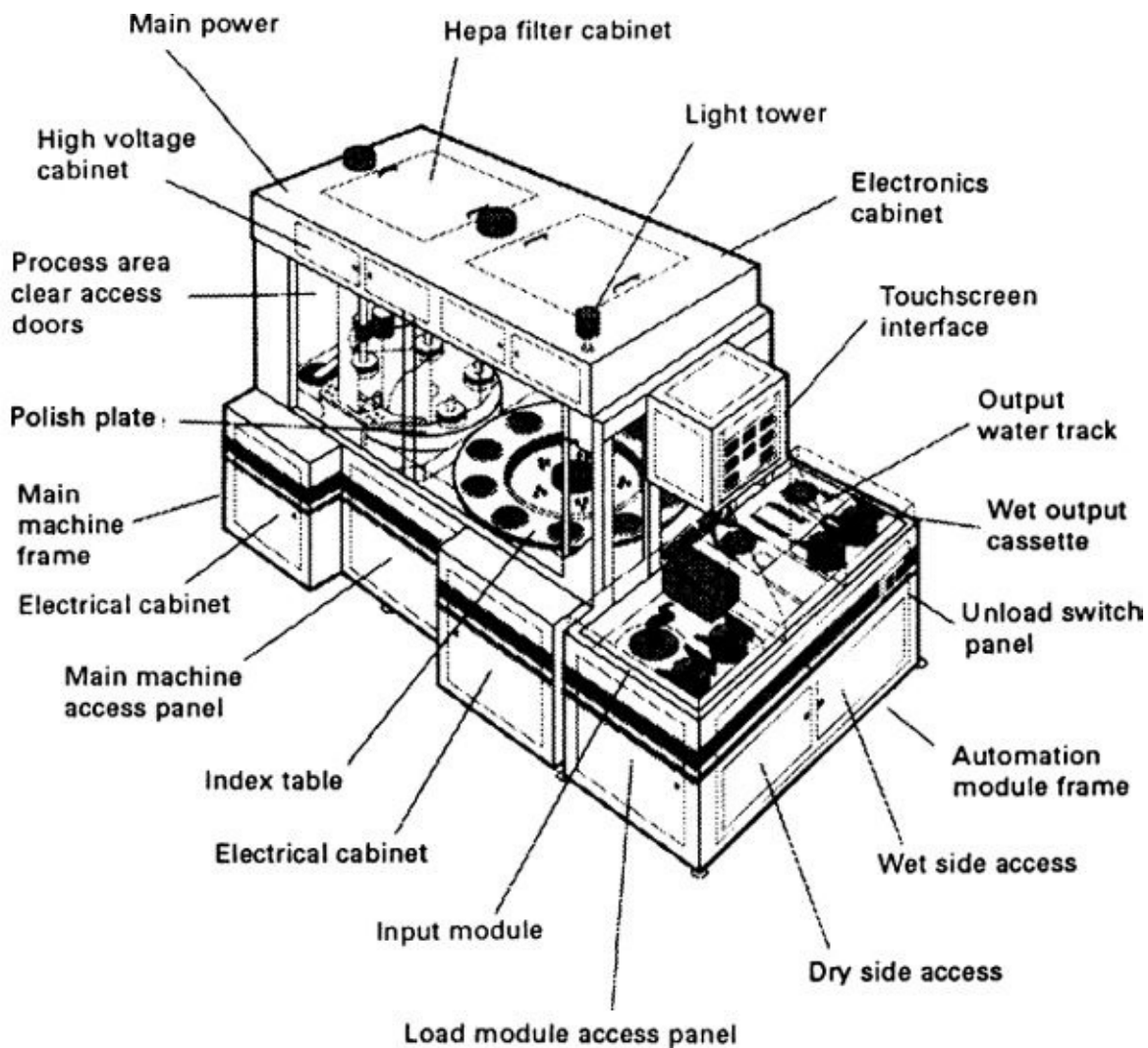


FIGURE 10.35 CMP system. (Source: SpeedFam CMP-V System, Semiconductor International, May 1993.)

CMP Summary

CMP is a critical planarization process that requires a high level of integration and balancing of many process parameters. The primary parameters are: polishing pad composition, polishing pad pressure, pad rotation speed, platen rotation speed, slurry flow rates, slurry chemistry, and slurry material selectivity. In addition to planarization for lithography improvements, CMP is the enabling process for dual-damascene patterning and copper metallization. This application is further explored in [Chap. 13](#).³³

Reflow

Some device schemes use a hard planarizing layer or layers. A popular layer is a deposited silicon dioxide doped with about 4 to 5 percent boron, called boron silicate glass (BSG). The presence of the boron causes the glass to flow at a relatively low temperature (less than 500°C), creating a planarized surface.

Another hard planarizing layer used is a spin-on-glass (SOG) layer. The glass is a mixture of silicon dioxide in a solvent that evaporates quickly. After spin application, the glass film is baked, leaving a planarized silicon dioxide film. The glass as spun is brittle, and some formulas contain between 1 and 10 percent carbon to increase resistance to cracking.

Image Reversal

The preference for positive resists over negative resists for small geometry patterning has been discussed. One of the advantages with positive resists is the use of dark-field masks for the imaging of holes. Dark-field masks offer a lower-defect process, because the majority of the surface is covered by hard chrome that does not damage like glass. However, some mask levels require the printing of islands rather than holes. Metal mask levels are island patterns. Unfortunately, the printing of an island with a positive resist requires the use of a clear-field mask with its glass damage potential.

A process that allows the printing of islands with positive resists and dark-field masks is image reversal, which involves the formation of the image in the resist with a dark-field mask by conventional masking steps ([Fig. 10.36](#)). At the conclusion of the exposure step, there is an image in the resist that is reversed from the desired image. That is, if the resist was developed, a hole rather than an island would be formed.

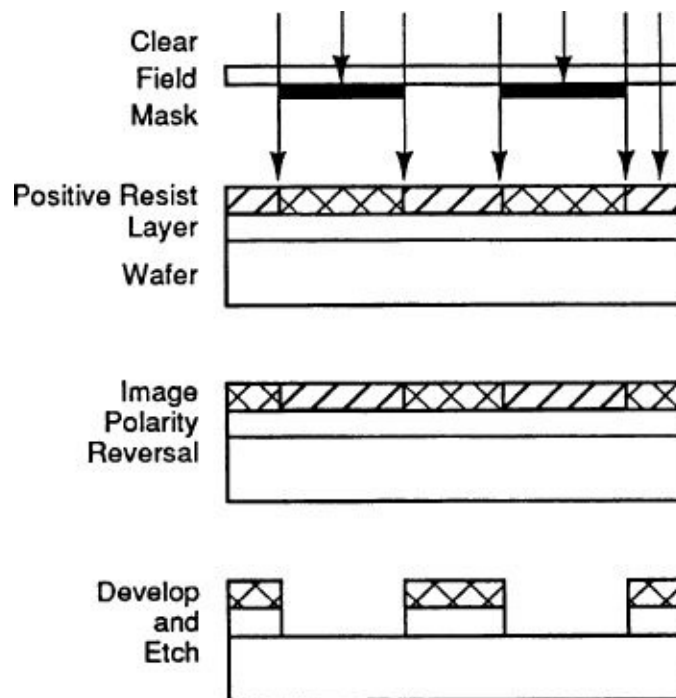


FIGURE 10.36 Image reversal.

The image reversal process involves exposing the resist-covered wafer to amine vapors in a vacuum oven. The vapors penetrate the resist, reversing its polarity. On removal of the wafers from the oven, they are given a flood exposure which completes the reversal process. The effect of the amine bake and flood exposure is

to change the relative dissolution rates of the exposed and unexposed regions, thus reversing the original image when the resist is finally developed. This process is capable of the same resolution capabilities as a nonreversed positive process.

Contrast Enhancement Layers

Optical projection system resolution is approaching the limits imposed by the constraints of the lens and the wavelength of the exposing radiation. The two set up a condition in which the resist *contrast threshold* becomes the limiting factor. This is because, at short exposure times and ultraviolet and deep ultraviolet energies, the energy of the exposing wave varies in intensity. Thus, the image formed in the resist is fuzzy.

A method used to decrease this threshold is a contrast enhancement layer (CEL), which is a layer spun on top of the resist that is initially opaque to the exposing radiation (Fig. 10.37).

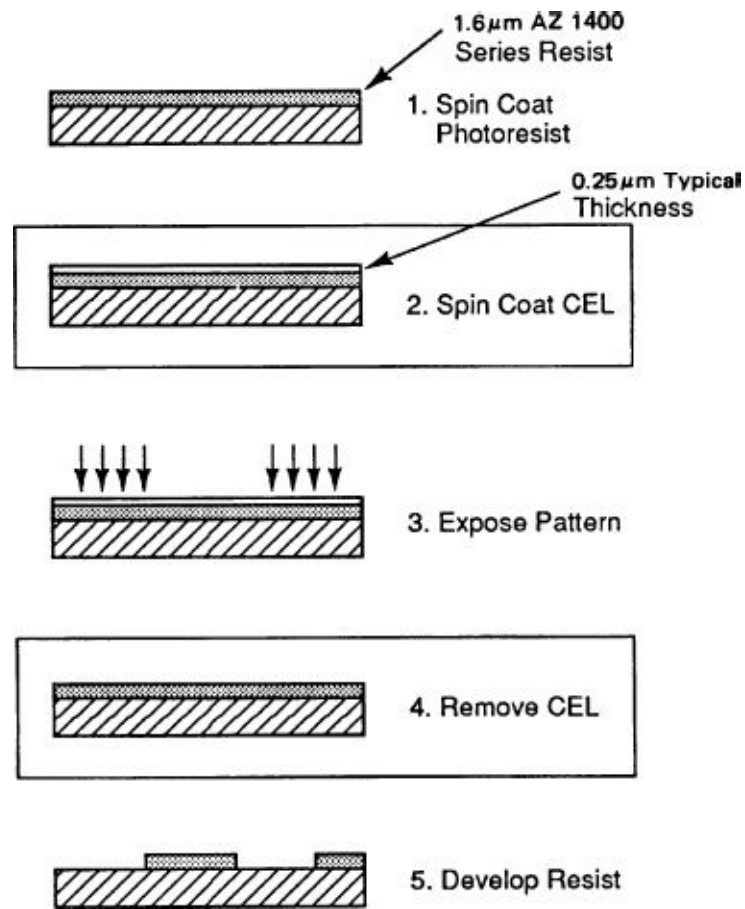


FIGURE 10.37 Contrast enhancement layer process flow.

During the exposure cycle, the CEL becomes bleached (transparent) and allows the radiation to pass into the underlying photoresist. The CEL responds first to the higher intensities before turning transparent, in effect storing the lower intensities before turning transparent. The result is that the resist receives a uniform exposure

of high-intensity radiation, which improves its resolution capability. Another way to imagine the role of the CEL is as the top layer of a dual-layer resist system with the image being formed in the thin top layer.

Before the resist is developed, the CEL is removed by a chemical spray and development of the resist proceeds by normal processing. Positive resist processes normally capable of 1.0- μm resolution can achieve a 0.5- μm image with a CEL.

Dyed Resists

Various dyes may be added into the resist during manufacture. A dye may have one or several effects during the exposure step. One possible effect is the absorption of radiation, thereby attenuating the reflected radiation and minimizing standing wave effects. Another is a change of the dissolution rate of the resist polymer during development. This effect creates a cleaner developed line (increased contrast).³⁵ An important use of dyes is the elimination of the notching that occurs in thin lines of deposited material crossing over surface steps. Addition of a dye to a resist can cause an increase in exposure time of 5 to 50 percent.³⁶

Improving Etch Definition

Forming the correct image in the photoresist is a critical step, but not the only step that defines the image in the wafer surface. Etching must also be controlled and precise. Several techniques are available that provide improved etch definition.

Lift-Off Process

The final dimensions of the images in the surface layer are the result of variations in both the exposure step and the etch step. In processes where etch undercutting (resist adhesion) is a problem, such as aluminum etching, the etch component of the dimensional variation can be the dominant one.

A patterning process that eliminates the etch variation component is lift-off ([Fig. 10.38](#)). In this variation, the wafer is processed through the development step, leaving a hole in the resist layer where a deposited layer is to be located. The exposure and development steps are adjusted to create a negative slope in the sidewall of the hole.

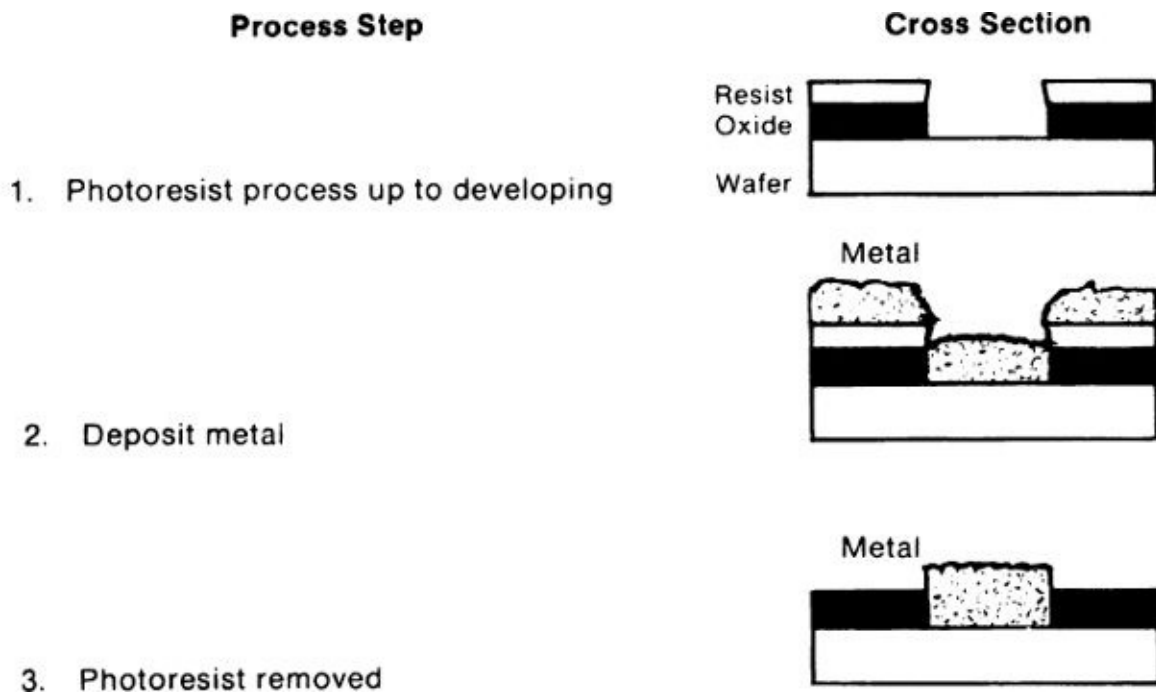


FIGURE 10.38 Lift-off.

Next, the wafer receives the deposited layer, which covers the entire wafer and fills in the hole. Definition of the pattern comes when the wafer is processed through a photoresist-removal step that lifts off the resist and unwanted metal layer.

Usually, the removal step is assisted by ultrasonic agitation. This helps form a clean break of the deposited film at the resist edge. After resist and film removal, the desired pattern is left on the wafer surface.

Self-Aligned Structures

Overetching has the effect of placing two structures closer together than intended. There is always some amount of overetching, and the alignment of some structures is absolutely critical. One solution is the *self-aligned structure*, such as an MOS gate ([Chap. 16](#)). The width and pattern integrity of the gate structure also defines the neighboring source or drain regions ([Fig. 10.39](#)). Opening the source or drain regions is a simple procedure of dip etching the oxide off the source or drain regions. The thinner oxide on the source or region allows a short etch time that does not allow enough time to etch the sides of the gate structure. The subsequent source or drain doping places the dopants next to the gate. This basic technique of a differential oxide thickness and dip etch is used to define or etch other structures. The gate structure functions as a doping block.

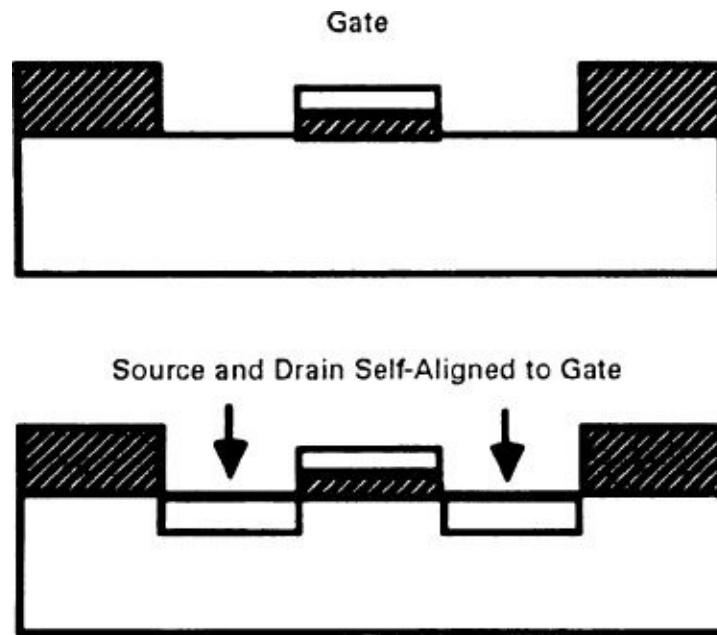


FIGURE 10.39 Self-aligned silicon gate (SAG) structure.

Etch Profile Control

In [Chap. 9](#), the concept of anisotropic etching was explained as a way of producing vertical (or near vertical) etched sidewalls. The problem becomes more complicated when the etched layer is actually a stack of different materials ([Fig. 10.40a](#)). Use of an etchant that has poor selectivity can produce a stack with varying widths ([Fig. 10.40b](#)). Profile control for multilayer etches becomes a trade-off of the etchant chemistry, power levels, and the system pressure and design.

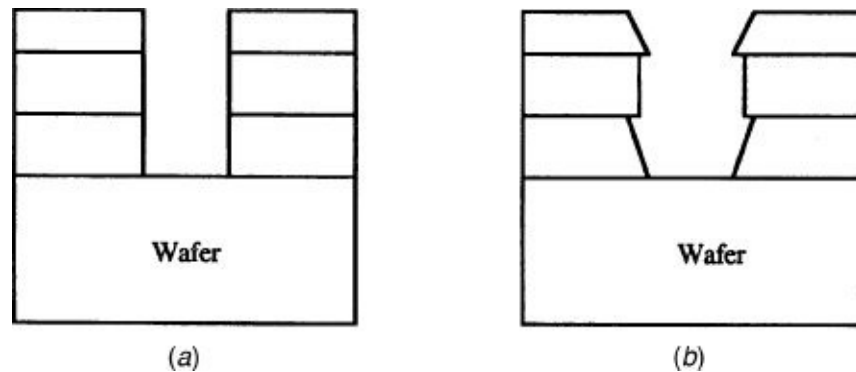


FIGURE 10.40 (a) Anisotropic and (b) isotropic etch of surface layer “stack.”

Review Topics

Upon completion of this chapter, you should be able to:

1. Describe four exposure-related effects that cause image distortion.
 2. Draw a cross-sectional flow diagram of a dual-layer resist process.
 3. Draw a cross-section flow diagram of a dual-damascene process.
 4. List two planarization techniques.
 5. State an advantage of an image-reversal process.
 6. Describe how antireflective layers, contrast enhancement layers, and resist dye additives improve resolution.
 7. Identify the parts of a pellicle and the advantages it offers to a resist process.
-

References

1. Mack, C., "Lithography, Forecast 1993, Fitting the Pieces Together," *Semiconductor International*, Cahners Publishing, Jan. 1993:31.
2. McCallum, M., Canning, J., and Shelden, G., "Lithography Trends," *Future Fab International* 9:145.
3. Staff, "Speeding the Transition to 0.018 μm ," *Semiconductor International*, Jan. 1998:66.
4. Executive Summary, SCALPEL Process, <http://www.lucnet.com>, May 2013.
5. Ghandhi, S. K., *VLSI Fabrication Principles*, 1994, John Wiley & Sons, Inc., New York, NY:687.
6. Zhang, Y., "Potential of KrF Scanning Lithography," *Future Fab International* 9:14.
7. Benschop, J., *EUV: Questions and Answers*, ASML Press, 2012, Oct. 15.
8. Elliott, D. J., *Integrated Circuit Fabrication Technology*, 1976, McGraw-Hill, New York, NY:82.
9. Zhang, Y., "Potential of KrF Scanning Lithography," *Future Fab International* 9:14
10. Ghandhi, S. K., *VLSI Fabrication Principles*, 1994, John Wiley & Sons, Inc., New York, NY:693.
11. Ultratech Stepper Inc., "Variable Numerical Aperture Large-Field Unit-Magnification Projection System," EP 1579259 A2, Sep. 28, 2005.
12. Nikon, *Immersion Lithography Technology Paper*, www.nikonprecision.com, April 2013.
13. Levenson, M. D., "Extending Optical Lithography to the Gigabit Era," *Microlithography World*, Autumn 1994:5.
14. Peters, L., "Reading Resists for the 90-nm Node," *Semiconductor International*, Feb. 2002:63.
15. Hand, A., "Intrinsic Birefringence Won't Halt 157-nm Lithography," *Semiconductor International*:42.
16. Reynolds, J., "Maskmaking Tour Video Course," Semiconductor Services, Redwood City, CA, Aug. 1991, Segment 10.
17. Spence, C., *Optical Proximity Photomask Manufacturing Issues*, OCG Microlithography Seminar Proceedings, 1984:255.
18. Ixcoff, R., "Pellicles 1985; An Update," *Semiconductor International*, Apr. 1985:111.
19. Micropel Division, *Micropel Product Data Sheet*, EKC Technology,

Hayward, CA, 1988.

[20.](#) Ling, C. H., and Liauw, K. L., “Improved DUV Multilayer Resist Process,” *Semiconductor International*, Nov. 1984:102.

[21.](#) Nishi, D., *Handbook of Semiconductor Manufacturing Technology*, 2nd ed., 2006, CRC Press, New York, NY:20–47.

[22.](#) Hand, A., “Intrinsic Birefringence Won’t Halt 157-nm Lithography,” *Semiconductor International*:42.

[23.](#) Moffatt, B., “Private Conversation,” *Yield Engineering Systems*, Livermore, CA, Jan. 17, 2013.

[24.](#) Steigerwald, J., Muraka, S., and Gutmann, R., *Chemical Mechanical Planarization of Microelectronic Materials*, John Wiley & Sons, Inc., Hoboken, NJ, 1997:4.

[25.](#) Peterson, M., Small, R., Shaw, G., et al., “Investigation CMP and Post-CMP Cleaning Issues for Dual-Damascene Copper Technology,” *Micro*, Jan. 1999:31.

[26.](#) Ibid.

[27.](#) Skidmore, K., “Techniques for Planarizing Device Topography,” *Semiconductor International*, Apr. 1988:116.

[28.](#) Ibid.

[29.](#) Jackson, R., Broadbent, E., Cacouris, T., et al., “Processing and Integration of Copper Interconnects,” *Solid State Technology*, Mar. 1998:49.

[30.](#) Ibid.

[31.](#) Ibid.

[32.](#) Iscoff, R., “CMP Takes a Global View,” *Semiconductor International*, Cahners Publishing, May 1994:74.

[33.](#) Steigerwald, J., Muraka, S., and Gutmann, R., *Chemical Mechanical Planarization of Microelectronic Materials*, John Wiley & Sons, Inc., Hoboken, NJ, 1997:4.

[34.](#) Housley, J., Williams, R., and Horiuchi, I., “Dyes in Photoresists: Today’s View,” *Semiconductor International*, Apr. 1988:142.

[35.](#) Elliott, D. J., *Integrated Circuit Fabrication Technology*, 1976, McGraw-Hill, New York, NY:168.

CHAPTER 11

Doping

Introduction

One of the unique properties of semiconductor materials is that their conductivity and the type of conductivity (N or P) can be created and controlled by introduction of specific dopants into the material. This concept was explored in the [Chaps. 2 and 3](#). A wafer starts into the wafer-fabrication process with either an N (negative) or P (positive) electrical conductivity. Through the fabrication process, the structures of the various transistors, diodes, resistors, and conductors are formed in and on the wafer surface. In this chapter, the formation of specific “pockets” of conductive regions and N-P in and on the wafer surface is described.

The structure that makes transistors and diodes function is an N-P or P-N junction. A *junction* is essentially the dividing line between a region that is rich in negative electrons (N-type region) and a region that is rich in holes (P-type region). The exact location of a junction is where the concentration of electrons equals the concentration of holes. This concept is illustrated later in the chapter.

Junctions are formed in the surface of a semiconductor wafer by the introduction of specific dopants (doping), by *ion implantation* or *thermal diffusion* processes. With thermal diffusion, dopant materials are introduced into the exposed top surface of the wafer, typically through a hole in the top silicon dioxide layer. With heating, they spread down into the bulk of the wafer. The amount and depth of the spread is governed by a set of rules, as explained below. These rules arise from a set of chemical rules that govern any movement of dopants in the wafer whenever the wafer is heated to a threshold temperature. In ion implantation, the dopant materials are literally shot into the wafer surface, with most of the dopant atoms coming to rest below the surface. Additional movements of the implanted atoms are also governed by the rules of diffusion ([Fig. 11.1](#)). Ion implantation has replaced the older thermal diffusion process for the introduction of dopants into wafers. Plus, ion implantation serves other roles in the fabrication of today’s miniaturized and multistructure devices. Thus, this chapter starts with a discussion of semiconductor junctions and concentration profiles, continues on with a brief overview of thermal diffusion processes, and ends with a description of the ion implant processes.

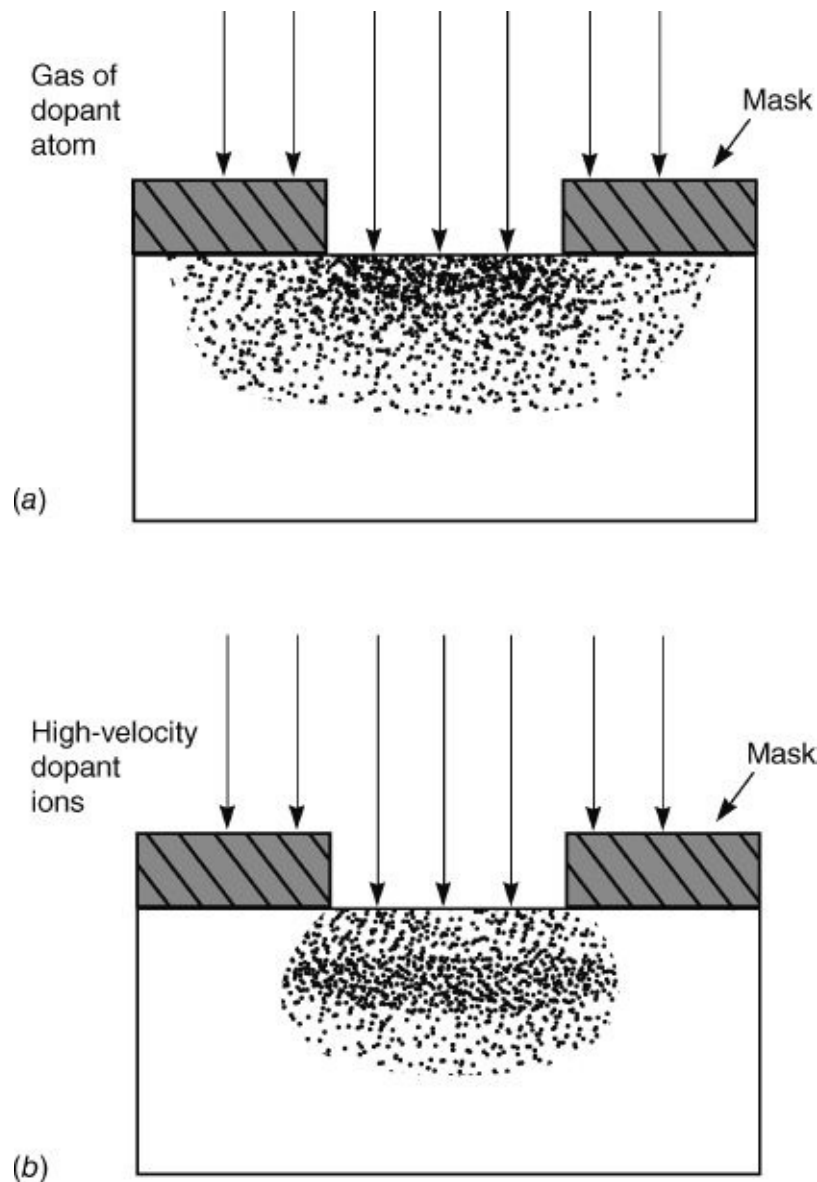


FIGURE 11.1 (a) Dopant concentration from diffusion and (b) dopant concentration from ion implantation.

The Diffusion Concept

A major advancement in semiconductor production was the development of diffusion doping. Diffusion, the movement of one material through another, is a natural chemical process with many examples in everyday life. Two conditions are necessary for a diffusion to take place. First, one of the materials must be at a higher concentration than the other. Second, there must be sufficient energy in the system for the higher-concentration material to move into or through the other. An example of gas-state diffusion is the action of a common pressurized spray can ([Fig. 11.2](#)), such as a room deodorant. When the nozzle is depressed, the pressurized material leaves the can and moves into the surrounding air. Thereafter, movement of the gas into the room proceeds by the process of diffusion. The movement ensues while the nozzle is depressed and continues after it is closed. The diffusion will continue as

long as the advancing spray is at a concentration higher than that in the air. As the material moves away from the can, the concentration of the material becomes progressively less. This is a characteristic of a diffusion process. Diffusion will continue until the concentration is even throughout the room.

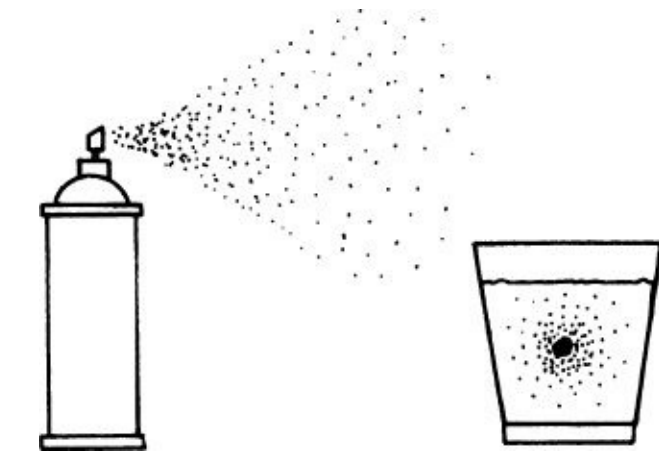


FIGURE 11.2 Examples of diffusion.

Another example is liquid-state diffusion as represented when a drop of ink is dropped into a glass of water. The ink is more concentrated than the surrounding water and immediately starts diffusing into the glass of water. The diffusion will continue until the whole glass of water is the same color. This example can be used to illustrate the influence of energy on the diffusion process. If the water in the glass is heated (giving the water more energy), the ink will spread through the water faster.

The same diffusion phenomenon takes place when a doped wafer is exposed to a concentration of atoms higher than the concentration in the wafer. This is called *solidstate* diffusion. These rules govern the movement of a dopant through the host wafer every time the wafer goes through a heat process with a temperature high enough to cause dopant movement, like during the anneal process after an ion implantation. Or, once a dopant is introduced into the wafer they will keep moving. Design rules must account for these movements and fabrication processes are often characterized by their “total thermal budget.”

Formation of a Doped Region and Junction

The formation of a doped region and junction is illustrated here by examining the doping of a wafer in a diffusion process, but the same concept applies to doping by ion implantation. The starting condition is depicted in [Fig. 11.3](#). The wafer illustrated is from a P-type crystal. The + symbols in the diagram represent the P-type dopants that were incorporated into the crystal during the crystal-growing

process. They are uniformly distributed throughout the wafer.

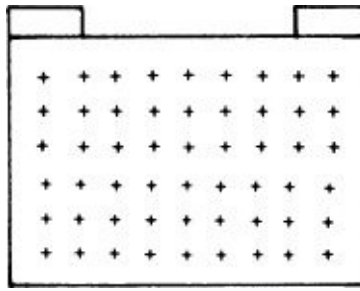


FIGURE 11.3 P-type wafer ready for diffusion.

The wafer receives a thermal oxidation and a patterning process that leaves a hole in the oxide layer. In a diffusion tube, the wafer is exposed to a concentration of N-type dopants at a high temperature (the + symbols in [Fig. 11.4](#)). The N-type dopants diffuse through the hole in the oxide layer.

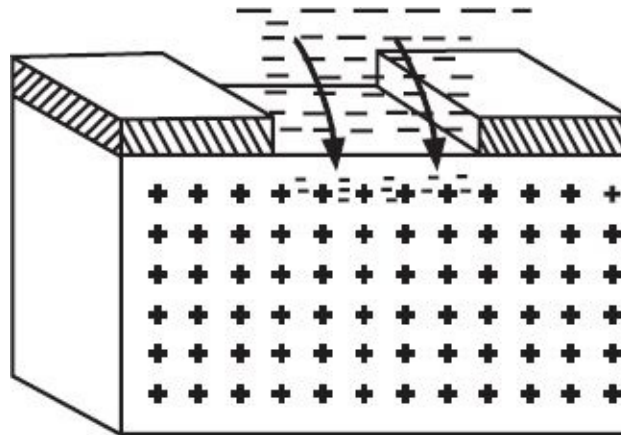


FIGURE 11.4 Start of a diffusion process.

The effect in the wafer is illustrated by examining what happens at different levels in the wafer. The conditions in the diffusion tube are set such that the number of N-type dopant atoms that diffuse into the wafer surface is greater than the number of P-type atoms in layer no. 1. In the illustration, there are seven more N-type atoms than P-type atoms, making that level electrically an N-type layer. In other words, this process has converted the top layer from P-type to N-type.

The diffusion process proceeds with N-type atoms diffusing from the first level down to the second level ([Fig. 11.5](#)). At the second level, there are again more N-type dopants than P-type dopants, converting level 2 to N-type. In the table ([Fig. 11.6](#)) is an accounting of the number of N-type and P-type atoms at each level. This process goes deeper into the wafer.

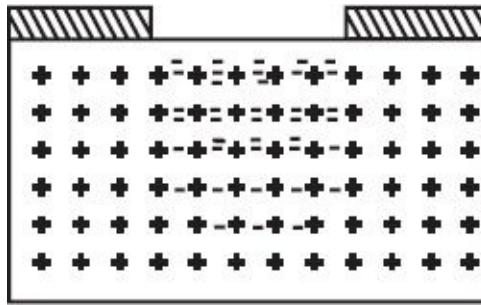


FIGURE 11.5 Cross-section of wafer at conclusion of diffusion.

Layer	# N's (-)	# P's (+)	Net (N - P)	Layer
1	12	5	7	N
2	10	5	5	N
3	8	5	3	N
4	5	5	0	Jct
5	3	5	-1	P
6	0	5	-5	P

FIGURE 11.6 Dopant amounts and level conductivity type.

The N-P Junction

At level 4, there are exactly the same number of N-type and P-type atoms. This level is the location of an N-P junction. The definition of an N-P junction is the location in the wafer where the numbers of N-type and P-type dopant atoms are equal. Note that below the junction at level 5, there are three N-type atoms, which are not enough to convert that layer to N-type.

The term *N-P junction* indicates that there is a higher concentration of N-type dopants on one side of the junction. A P-N junction would indicate more P-type dopants on one side of the junction.

The behavior of electrical currents across semiconductor junctions gives rise to the particular performance of individual semiconductor devices and is the subject of [Chap. 14](#). In this chapter, the emphasis is on the formation and character of doped regions in a wafer.

Doping Process Goals

The goals of a doping process, whether it is ion implantation or thermal diffusion, are threefold:

1. To create a specific number (concentration) of dopant atoms at and below the wafer surface
2. To create an N-P or P-N junction at a specific distance below the wafer surface
3. To create a specific distribution and concentration

of dopant atoms in the wafer surface **Graphical Representation of Junctions**

In a cross-section of a semiconductor device ([Fig. 11.9](#)), the N-P junctions are indicated simply as regions in the device. There is no graphical convention to represent N-or P-type areas. The drawings just show the relative location of the doped region and the junction. This type of graphical representation gives little information about the concentration of the dopant atoms and only approximates the actual dimensions of the regions. A drawing of a 2- μm -deep junction in a 20-mm-thick wafer, scaled to an 8-ft wafer, would have a junction depth only 0.4 in thick.

Concentration versus Depth Graphs

Another two-dimensional graphical representation of a doped region is the concentration-versus-depth graph. The concentration is represented on the vertical axis and the depth into the wafer on the horizontal axis. An example of such a graph is illustrated in [Fig. 11.7](#). The illustration uses the data from the doping example shown in [Fig. 11.6](#). First, the P-type dopant concentration is plotted. In the example, there are exactly five P-type dopant atoms at every level, resulting in a straight horizontal line on the graph ([Fig. 11.7b](#)). Next, the number of N-type dopant atoms is plotted. Since the number of atoms decreases deeper into the wafer, the plotted line slopes down and to the right. At level 4, the number of N-and P-type dopants is equal, and the lines cross. This is a graphical representation of the location of the junction.

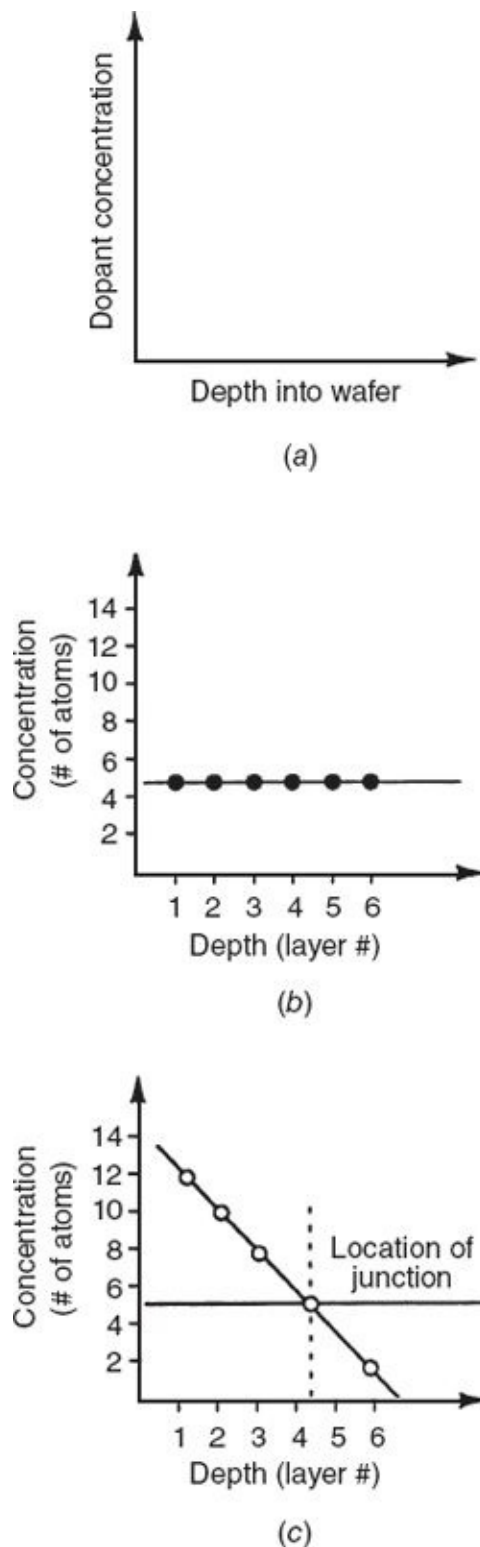


FIGURE 11.7 Construction of concentration-versus-depth curve: (a) axes, (b) P-type dopant, and (c) N-type and P-type dopant.

The graph of an incoming dopant concentration versus depth profile for an actual process is not a straight line. They are curved lines. The shape of the curve is determined by the physics of the dopant technique. The actual shapes are discussed

in the deposition and drive-in sections.

Lateral Diffusion

The diffusion doping process depicted in [Fig. 11.5](#) shows the incoming dopant atoms traveling straight down into the wafer. In reality, the dopant atoms move in all directions. An accurate cross-section ([Fig. 11.8](#)) would show that some of the atoms have moved in a lateral direction, forming a junction under the oxide barrier. This movement is also called *lateral* or *side diffusion*. The amount of the lateral or side diffusion is approximately 85 percent of the vertical junction depth. Lateral diffusion takes place regardless of whether the introduction was through diffusion or ion implantation. The effects of side diffusion on circuit density are discussed in the introduction to ion implantation.

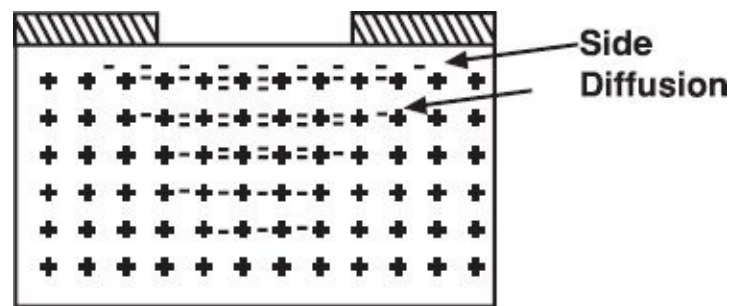


FIGURE 11.8 Side-diffused N-type dopants.

Same-Type Doping

Some devices call for a doping with a dopant type the same as the host region. In other words, an N-type dopant will be put into an N-type wafer, or a P-type dopant will be put into a P-type wafer. When this situation happens, the added dopant atoms simply increase the concentration of the dopant atoms in the localized region. No junction is formed.

Diffusion Process Steps

The use of *solidstate thermal diffusion* to create junctions in semiconductor wafers requires two steps. Step one is called *deposition*, and step two is called *drive-in oxidation*. Both steps take place in a horizontal or vertical tube furnace. The equipment is the same as described in [Chap. 7](#) for oxidation.

Diffusion Step Purpose

1. Deposition Introduce dopant(s) into wafer surface

2. Drive-In Drive (spread) the dopants to desired depth

Deposition

Deposition (also called *predeposition*, *dep*, or *predep*) takes place in a tube furnace, with the wafers placed on a quartz “boat” in the flat zone of the tube. A source of dopant atoms is located in the source cabinet and their vapors are transferred into the tube at a required concentration ([Fig. 11.9](#)). Liquid, gas, and solid dopant sources are used.

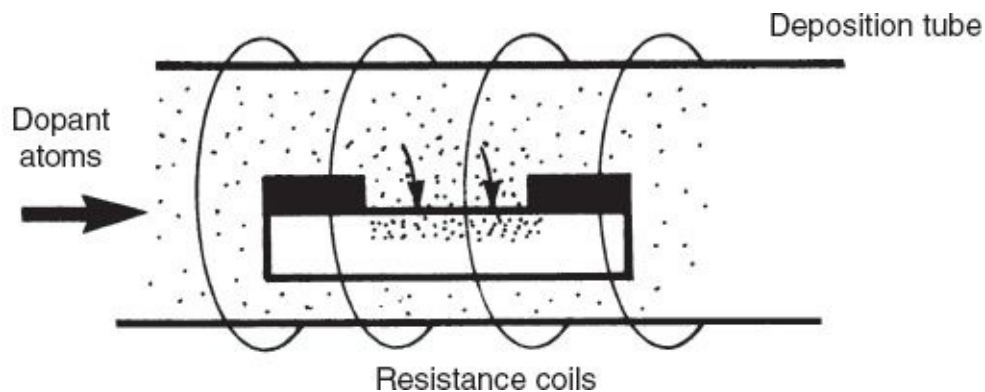


FIGURE 11.9 Deposition.

In the tube, the dopant atoms diffuse into the exposed wafer. Within the wafer, the dopant atoms move by two different mechanisms: vacancy and interstitial movement. In the vacancy model ([Fig. 11.10a](#)), the dopant atoms move by filling empty crystal positions, called *vacancies*. The second model ([Fig. 11.10b](#)) relies on interstitial movement of the dopant.¹ In this model, the dopant atom moves through the spaces between the crystal sites, that is, interstitial.

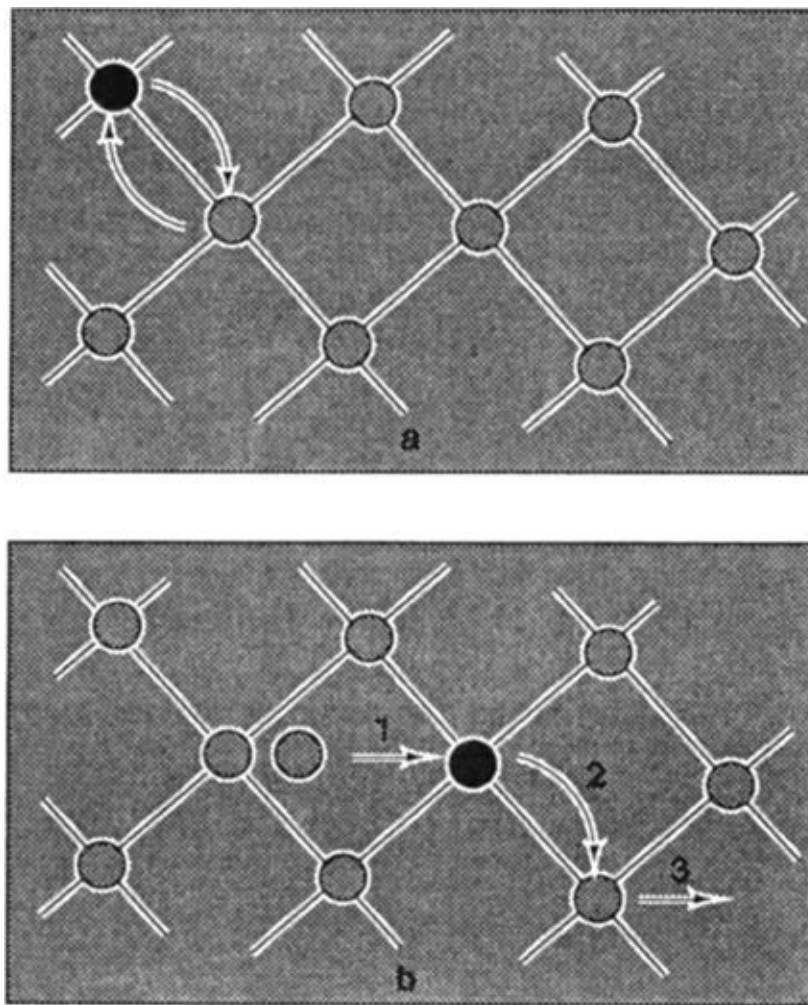


FIGURE 11.10 Diffusion models: (a) vacancy model and (b) interstitial model.

A deposition process is controlled or limited by several factors. One is the *diffusivity* of the particular dopant. Diffusivity is the rate (speed) of movement of the dopant through the particular wafer material. The higher the diffusivity, the faster the dopant moves through the wafer. Diffusivity increases with temperature.

Another factor is the *maximum solid solubility* of the dopant in the wafer material. It is the maximum concentration of a specific dopant that can be put into the wafer. A familiar analogy is the maximum liquid solubility of sugar in coffee. The coffee can dissolve only a certain amount of sugar before it collects in the bottom of the cup as a solid.

The maximum solid solubility limit increases with increasing temperature. In a semiconductor deposition step, the concentration of the dopant is purposely set higher than the maximum solid solubility of the dopant in the wafer material. This situation ensures that the maximum amount of the dopant will be accepted by the wafer.

The amount of dopant entering the wafer surface is a function of the temperature only, and the deposition is said to take place at solid solubility. Maximum solubility

levels for various dopants in silicon are shown in [Fig. 11.11](#).

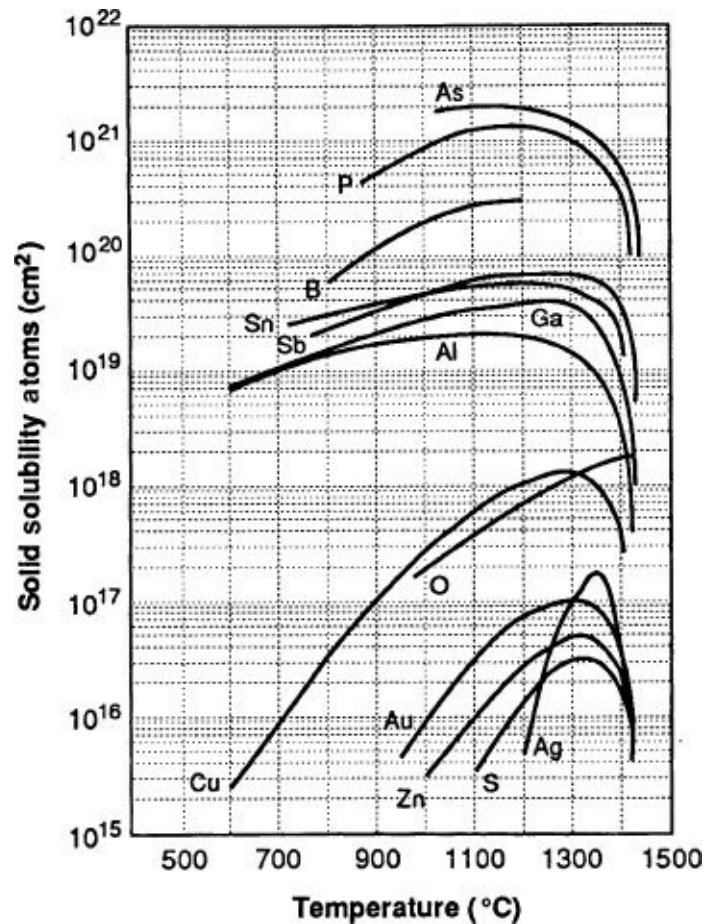


FIGURE 11.11 The solid solubility of impurities in silicon.

The concentration of dopant atoms at each level in the wafer is an important factor in the performance of junction diodes and transistors. A dopant concentration versus depth curve for a deposition is shown in [Fig. 11.12](#). The shape of the curve follows a mathematical formula specifically known as an *error function*. An important factor in device performance is the concentration of the dopant at the wafer surface. This is called the *surface concentration* and is the quantity indicated where the error function curve intersects the vertical axis. Another deposition parameter is the total quantity of atoms diffused into the wafer. This amount increases with the time of the deposition. Mathematically, the quantity of atoms (Q) is represented by the area under the error function curve.

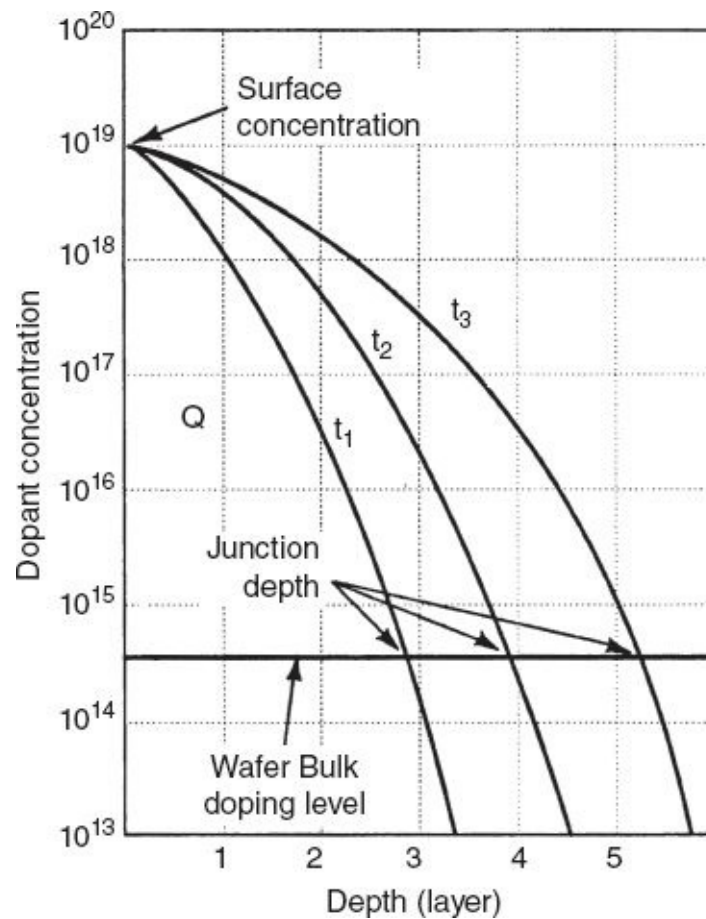


FIGURE 11.12 Typical deposition (error function) dopant profile for three different deposition times.

Lateral Diffusion

The diffusion doping process depicted shows the incoming dopant atoms traveling straight down into the wafer. In reality, the dopant atoms move in all directions. An accurate cross-section ([Fig. 11.8](#)) would show that some of the atoms have moved in a lateral direction, forming a junction under the oxide barrier. This movement is also called *lateral* or *side diffusion*. Lateral diffusion takes place regardless of whether the introduction was through diffusion or ion implantation. The effects of side diffusion on circuit density are discussed in the introduction to ion implantation.

Same-Type Doping

Some devices call for a doping with a dopant type the same as the host region. In other words, an N-type dopant will be put into an N-type wafer, or a P-type dopant will be put into a P-type wafer. When this situation happens, the added dopant atoms simply increase the concentration of the dopant atoms in the localized region. No junction is formed.

Dopant Sources

A deposition depends on the presence of a concentration of dopant atom vapors in the tube. The vapors are created from a dopant source located in the source cabinet of the tube furnace and passed into the tube with a carrier gas. Dopant sources are either in liquid, gaseous, or solid states. Several dopant elements are available in more than one state ([Fig. 11.13](#)).²

Type	Element	Compound Name	Formula	State	Diffusion Reactions*
N	Antimony	Antimony Trioxide	Sb ₂ O ₃	Solid	
	Arsenic	Arsenic Trioxide	As ₂ O ₃	Solid	2AsH ₃ + 3O ₂ → As ₂ O ₃ + 3H ₂ O
		Arsine	AsH ₃	Gas	
	Phosphorus	Phosphorus Oxychloride	POCl ₃	Liquid	4POCl ₃ + 3O ₂ → 2P ₂ O ₅ + 6Cl ₂
		Phosphorus Pentoxide	P ₂ O ₅	Solid	2PH ₃ + 4O ₂ → P ₂ O ₅ + 3H ₂ O
		Phosphine	PH ₃	Gas	
P	Boron	Boron Tribromide	BBr ₃	Liquid	4BBr ₃ + 3O ₂ → 2B ₂ O ₃ + 6Br ₂
		Boron Trioxide	B ₂ O ₃	Solid	B ₂ H ₆ + 3O ₂ → B ₂ O ₃ + 3H ₂ O
		Diborane	B ₂ H ₆	Gas	
		Boron Trichloride	BCl ₃	Gas	BCl ₃ + 3H ₂ → 2B + 6HCl
		Boron Nitride	BN	Solid	
	Gold	Gold	Au	Solid (Evap.)	
	Iron		Fe		
	Copper		Cu		
	Lithium		Li		Undesirable impurities from contamination
	Zinc		Zn		
	Manganese		Mn		
	Nickel		Ni		
	Sodium		Na		

*Note: Only selected diffusion reactions are listed

FIGURE 11.13 Deposition source table.

Liquid dopants are metered into the deposition tube from quartz flasks (bubblers) by an inert carrier gas ([Fig. 11.14](#)). Gas dopants are metered into the tube from pressurized tanks through a manifold ([Fig. 11.15](#)).

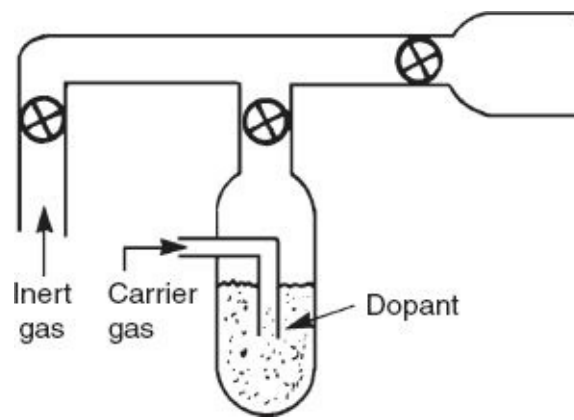


FIGURE 11.14 Liquid dopant source.

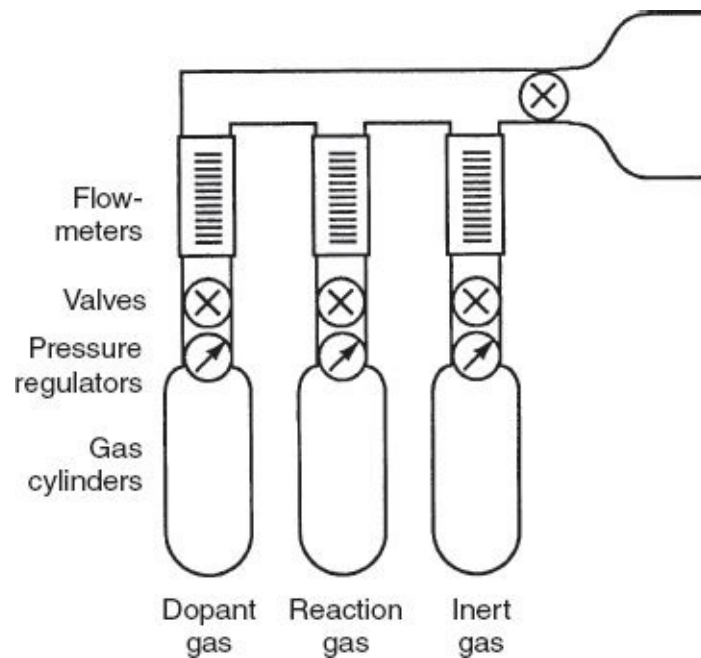


FIGURE 11.15 Gas source manifold.

A popular solid source is the planar source wafer. These are wafer-size “slugs” that contain the desired dopant. Boron slugs are a compound of boron and a nitride (BN). Slugs are also available for arsenic and phosphorus diffusions.

The slugs are stacked on the deposition boat, with one slug between every two device wafers. This arrangement is called a *solid neighbor source*. In the tube, the dopant diffuses out of the slug, crosses the short distance to the wafer, and diffuses into the surface.

The third solid dopant source is a conformal layer spun directly on the wafer surface. The sources are powdered oxides (same as remote sources) mixed in solvents. Left on the surface is a layer of doped oxide that conforms to the wafer surface. The heat of the deposition furnace drives the dopant out of the oxide and into the wafer.

Drive-In Oxidation

The second major part of the diffusion process is the drive-in-oxidation step. It is also variously known as *drive-in*, *diffusion*, *reoxidation*, and *reox*. The purpose of this step is twofold: redistribution of the dopant in the wafer and growth of a new oxide on the exposed silicon surface.

1. The first step is redistribution of the dopant deeper into the wafer. During the deposition, a high-concentration layer and a shallow layer of dopant are diffused into the surface. In the drive-in, there is no dopant source. The heat alone drives the dopant atoms deeper and wider into the wafer just as material from a spray can continue to spread into the room after the nozzle is released. During this step, the total amount of atoms (Q) from the deposition step remains constant. The surface concentration is reduced, and the distribution of atoms takes a new shape. The distribution after the drive-in is described by mathematicians as a Gaussian distribution ([Fig. 11.16](#)). The junction depth increases. Generally, the drive-inoxidation step takes place at a higher temperature than the deposition step.

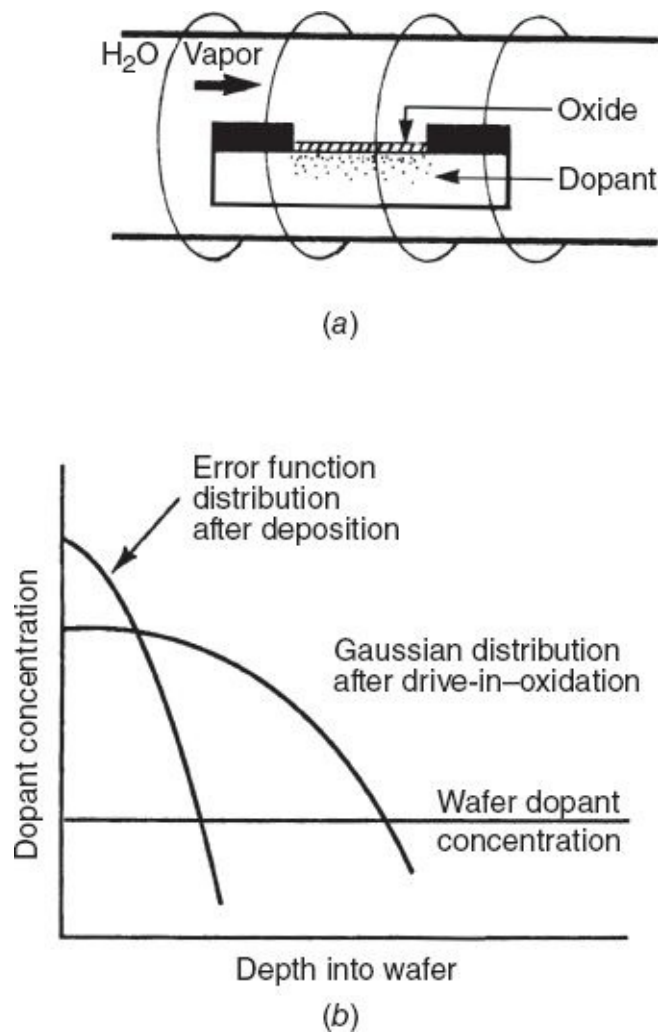


FIGURE 11.16 Drive-in oxidation: (a) cross-section of wafer and (b) dopant concentration in wafer.

2. The second purpose of drive-in oxidation is the oxidation of the exposed silicon surfaces. The atmosphere in the tube is oxygen or water vapor, which performs the oxidation simultaneously as the dopants are being driven deeper into the wafer.

The setup, process steps, and equipment for the drive-in-oxidation step are the same as an oxidation process.

Oxidation Effects

The oxidation of the silicon surface affects the final distribution of the dopants.³ The effects are related to the relocation of the top-level dopants after the oxidation. Recall that the silicon in the silicon dioxide film is consumed from the wafer surface. The question to ask is, “What happened to the dopants that were in the top level?” The answer to that question depends on the conductivity type of the dopant.

If the dopant is an N-type, an effect called *pile-up* ([Fig. 11.17a](#)) occurs. As the

oxide/silicon interface advances into the surface, the N-type dopant atoms segregate into the silicon rather than the oxide. The effect is to increase the number of these dopants in the new top layer of the silicon. In other words, the N-type dopants pile up in the wafer surface, and the surface concentration of the dopant is increased. Pile-up changes the electrical performance of the devices.

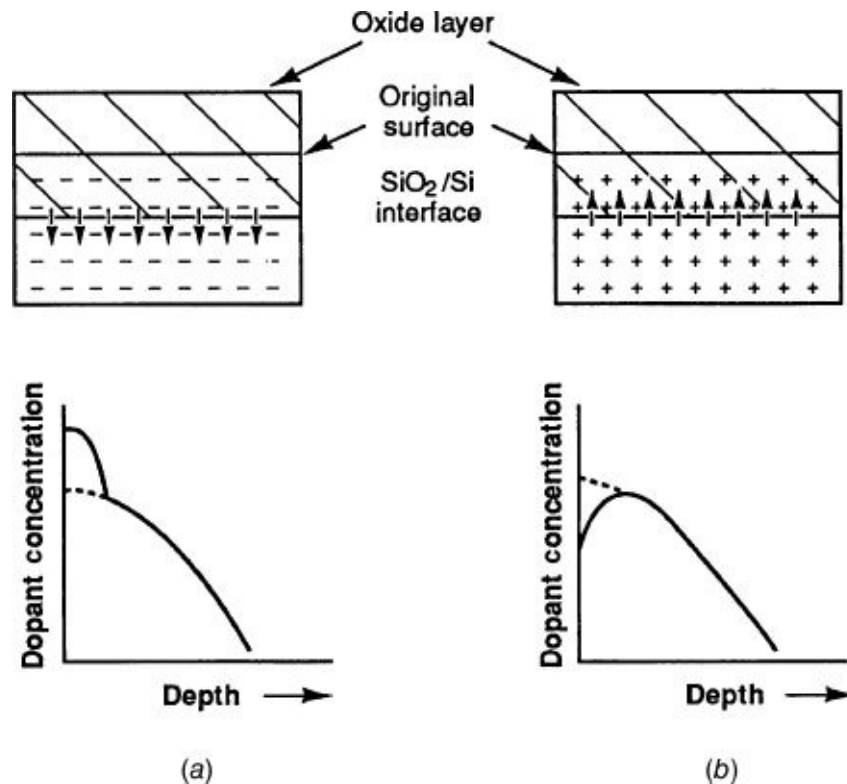


FIGURE 11.17 Pile-up and depletion of dopants during oxidation: (a) pile-up of N-type dopants and (b) depletion of P-type dopants.

If the dopant is the P-type boron, an opposite effect occurs. The boron atoms are more soluble in the oxide and are drawn up into it ([Fig. 11.17b](#)). The effect on the wafer surface is a lowering of the concentration of boron atoms, which lowers the surface concentration and also affects the electrical performance of the devices. A summary of the deposition and drive-in-oxidation steps is provided in [Fig. 11.18](#).

	Deposition	Drive-In
Goals	Introduction of Dopant	1. Redistribution of Dopant 2. Reoxidation
Variables		1. Surface Composition 2. Junction Depth 3. Time 4. Diffusivity 5. Temperature 6. Quantity of Atoms
Source Conditions	Continuous Source	No Source
Temperature Range	900 - 1100°C	1050 - 1200°C
Oxidation	No	Yes

FIGURE 11.18 Summary of deposition and drive-in steps.

Introduction to Ion Implantation

Thermal diffusion places a limit on the production of advanced circuits. Five challenges are lateral diffusion, ultra-thin junctions, poor doping control, surface contamination interference, and dislocation generation. Lateral diffusion occurs during deposition and drive-in but also continues every time the wafer is heated into a range where diffusion movement can take place ([Fig. 11.19](#)). The circuit designer must leave enough room between adjacent regions to prevent the laterally diffused regions from touching and shorting. The accumulative effect for a dense circuit can be a largely increased die area. Another problem with high-temperature processing is crystal damage. Every time a wafer is heated and cooled, crystal damage from dislocations occurs. A high concentration of these dislocations can cause device failure from leakage currents. One goal of an advanced process sequence is a reduced *thermal budget* to reduce these two problems.

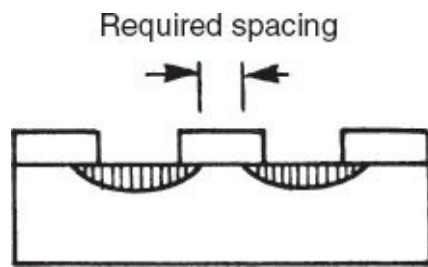


FIGURE 11.19 Side diffusion.

The advent of MOS transistors created two new doping requirements: low dopant concentration control and ultra-thin junctions. Gate regions with dopant concentrations below 10^{15} atom/cm² are required for efficient MOS transistors. However, this level is difficult to achieve consistently with a diffusion process. Scaling MOS transistors smaller to achieve higher packing densities also requires thin junction depths in the source/drain areas.⁴ Junction depths have continued to fall with sub-10-nm junctions projected in 2016.⁵

A fourth problem is imposed by the physics or mathematics of a diffused region. As illustrated in [Fig. 11.12](#), the majority of the dopant atoms are located near the wafer surface. This puts the majority of the current traveling near the surface where the dopant atoms are located. Unfortunately, this is the same location as contaminants (in and on the wafer surface) that interfere or degrade the current flow. A requirement of advanced devices is special wells in the surface with specific dopant gradients that cannot be achieved with diffusion technology. These wells allow high-performance transistors (see [Chap. 16](#)).

Ion implantation overcomes these limits of diffusion and also adds additional

benefits. Ironically, while ion implantation is the modern doping process, the technology has a long history. Machines were built in the 1940–1950s based on early work done by physicist Robert Van de Graff at MIT and Princeton. In 1954, William Shockely (yes, that Shockely) filed a patent for the use of ion implantation for the manufacture of semiconductors.⁶

During the ion-implant process, there is no side diffusion; the process takes place at close to room temperature; the dopant atoms are placed below the wafer surface; and a wide range of doping concentrations are possible. With ion implantation, there is greater control of the location and number of dopants put in the wafer. Also, photoresist and thin metal layers can be used as doping barriers along with the usual silicon dioxide layers. Given the benefits, it is not surprising that all of the doping steps for advanced circuits are done by ion implantation.

Concept of Ion Implantation

Diffusion is a chemical process. Ion implantation is a physical process, that is, the act of implanting does not rely on a chemical interaction between the dopant and the wafer material. An analogy that demonstrates the concept of ion implantation is a cannon firing balls into a wall ([Fig. 11.20](#)). Given enough momentum from the powder in the cannon, the balls will penetrate the wall and come to rest below the surface of the wall. The same events take place in an ion-implantation machine. Instead of cannonballs, dopant atoms are ionized and isolated, accelerated (gain momentum), formed into a beam, and swept across the wafer. The dopant atoms physically bombard the wafer, enter the surface, and come to rest below the surface ([Fig. 11.21a](#)).



FIGURE 11.20 Ion-implantation analogy.

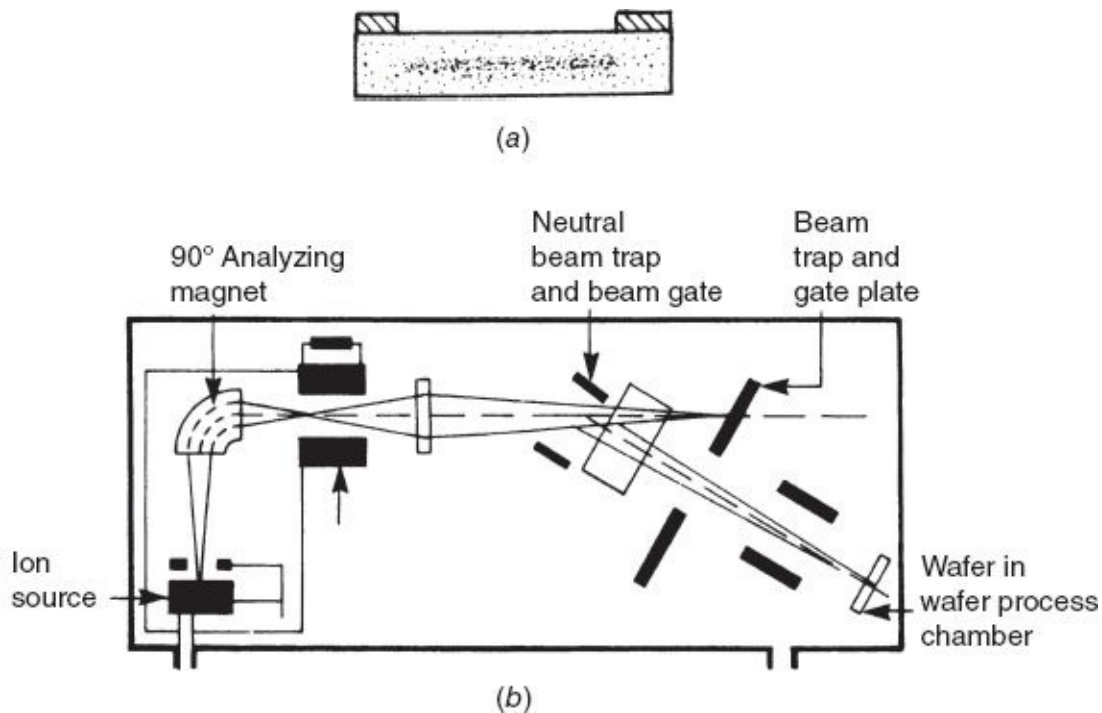


FIGURE 11.21 Ion implantation: (a) distribution of implanted atoms in wafers and (b) block diagram of ion implanter.

Ion-Implantation System

An ion implanter is a collection of very sophisticated subsystems ([Fig. 11.21b](#)); each performing a specific action on the ions. Ion implanters come in a variety of designs used for advanced research and/or high-volume production. All of the machines have the same major subsystems as described below.

Production-level ion implanters are designed to achieve the following:

- Accommodate a wide number of dopant species
- Implant uniformity across the wafer, wafer to wafer, and lot to lot
- Low-contamination levels
- Meet productivity levels

Implant Species Sources

The same dopant elements used in diffusion processes are used in ion-implant processes. In diffusion processes, the dopants, or species, originate gas, or solid sources.

Gases are favored for ion implantation because of their ease of use and higher control. The gases most used are AsH_3 , PH_3 , and BF_3 . One of the advantages of ion implant is a wider range of materials. Silicon (SiF_2) and germanium (GeF_4) can be implanted. Elemental arsenic and elemental phosphorus are solid sources used in

implanters. The gas cylinders are connected to the ion-source subsystem through mass-flow meters, which offer more control of the gas flow than normal flowmeters.

Ionization Chamber

The name “ion implant” implies that ions are a part of the process. Recall that ions are atoms or molecules with a negative or positive charge. The ions implanted are ionized atoms of the dopants. The ionization occurs in a chamber that is fed the source vapors. The chamber is maintained at a low pressure (vacuum) of about 10^{-3} torr. Inside the chamber is a filament that is heated to the point where electrons are created from the filament surface. The negatively charged electrons are attracted to an oppositely charged anode in the chamber. During the travel from the filament to the anode, the electrons collide with the dopant source molecules and create a host of positively charged ions from the elements in the molecule. The results of the ionization of the source BF_3 are shown in [Fig. 11.22](#).

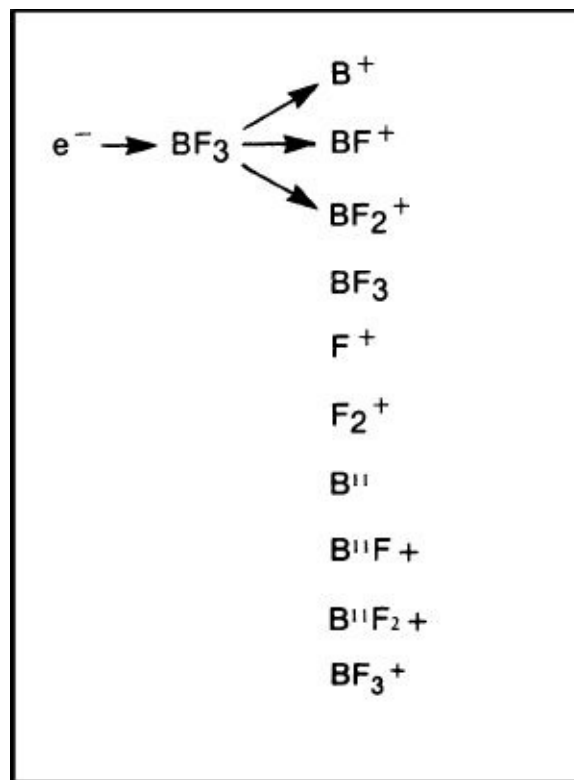


FIGURE 11.22 Ion species of BF_3 .

Another ionization method uses a cold-cathode technique to generate the electrons. A high-voltage electric field is created between a cathode and anode, which creates the electrons in a self-sustaining process.

Mass Analyzing or Ion Selection

At the top of the list in [Fig. 11.22](#) is a lone boron ion. This is the atom that is desired in the wafer surface. The other species resulting from the ionization of the boron trifluoride are not wanted in the wafer. The boron ion must be *selected* from the group of positive ions. This process is called *analyzing, mass analyzing, selection, or ion separation*.

Selection is accomplished in a mass analyzer. This subsystem was first developed during the Manhattan Project for the atomic bomb. The analyzer ([Fig. 11.23](#)) creates a magnetic field. The species leave the ionization subsystem with voltages of 15 to 40 keV (thousand electron volts). In other words, they are traveling at a relatively high speed.

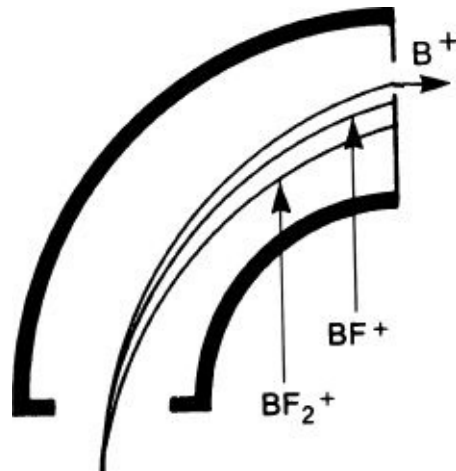


FIGURE 11.23 Analyzing magnet.

In the field, each of the positively charged species is bent into an arc with a specific radius. The radius of the arc is dictated by the mass of the individual species, its speed, and the strength of the magnetic field. At the exit end of the analyzer is a slit that will allow only one species to exit. The magnetic field is adjusted to match the path of the boron ion to the exit slit position. Thus, only the boron ion leaves the analyzing subsystem.

In some systems, the analysis also takes place after acceleration ([Fig. 11.23](#)). After acceleration, analysis is necessary if the implant species is a molecule that might separate in the acceleration process and/or to ensure an uncontaminated beam.

In some cases, the family of species separated in the analyzer contain some components that are closer to the mass of the intended implant species. These are called *mass interference*. They cannot be resolved in the analyzing magnet and end up in the implanting beam. Otherwise, there can be atoms of the intended species that have different energies yet the same magnetic characteristics. These also can end up in the implanting beam and eventually in the wafer.²

Acceleration Tube

On leaving the analyzing section, the boron ion moves into an acceleration tube. The purpose is to accelerate the ion to a high-enough velocity, thus gaining sufficient momentum to penetrate the wafer surface. *Momentum* is defined as the product of the mass of the atom multiplied by its velocity. This section is kept at a high vacuum (low pressure) to minimize contaminants from entering into the beam. Turbo-vacuum pumps ([Chap. 13](#)) are typically used for this purpose.

The required velocity is achieved by taking advantage of the fact that negative and positive charges attract each other. The tube is a linear design with annular anodes along its axis. Each of the anodes has a negative charge. The charge amount increases down the tube.

As the positively charged ion enters the tube, it immediately starts to accelerate down the tube. The voltage value is selected based on the mass of the ion and the momentum required at the wafer end of the implanter. The higher the voltage, the higher the momentum, and the faster and deeper the dopant ion can be implanted. Voltages range from 5 to 10 keV for low-energy implanters to 0.2 to 2.5 MeV (million electron) for high-energy implanters.

Ion implanters are classified into the categories of medium-and high-current machines, high-energy devices, and oxygen-ion implanters. The stream of positive ions exiting the tube is actually an electric current. The beam current level translates into the number of ions implanted per minute. The higher the current, the more atoms are implanted. The amount of atoms implanted is called the *dose*. Medium-current machines produce currents in the 0.5 to 1.7 mA (milliampere) range at energies from 30 to 200 keV (thousand electron volts). High-current machines generate beam currents of about 10 mA at energies up to 200 keV.⁹ High-energy implanters are finding use in several CMOS doping applications, including retrograde wells, channel stops, and deep-buried layers (see [Chap. 16](#)).

Wafer Charging

High-current implants create an unacceptable degree of electrical charging (*wafer charging*) of the wafer surface. The high-current beam carries excess positive charges that charge the wafer surface. The positive charge draws neutralizing electrons from the surface, the bulk, and from the beam. High-charge levels can degrade or destroy surface dielectric layers. Wafer charging is a particular problem on thinner MOS gate dielectrics.¹⁰ Methods used to neutralize or reduce the charge are *flood guns* specifically designed to provide electrons, a plasma bridge method of providing low-energy electrons,¹¹ and control of the electron path with magnetic fields.¹²

The relationships of current and energy for production-level implanters is shown

in [Fig. 11.24](#). High-energy machines accelerate ions in the 10 keV to 3.0 MeV range, with beam currents up to 1.0 mA. Oxygen-implant machines are high-current implanters used to implant oxygen in SOI applications (see [Chap. 16](#)).

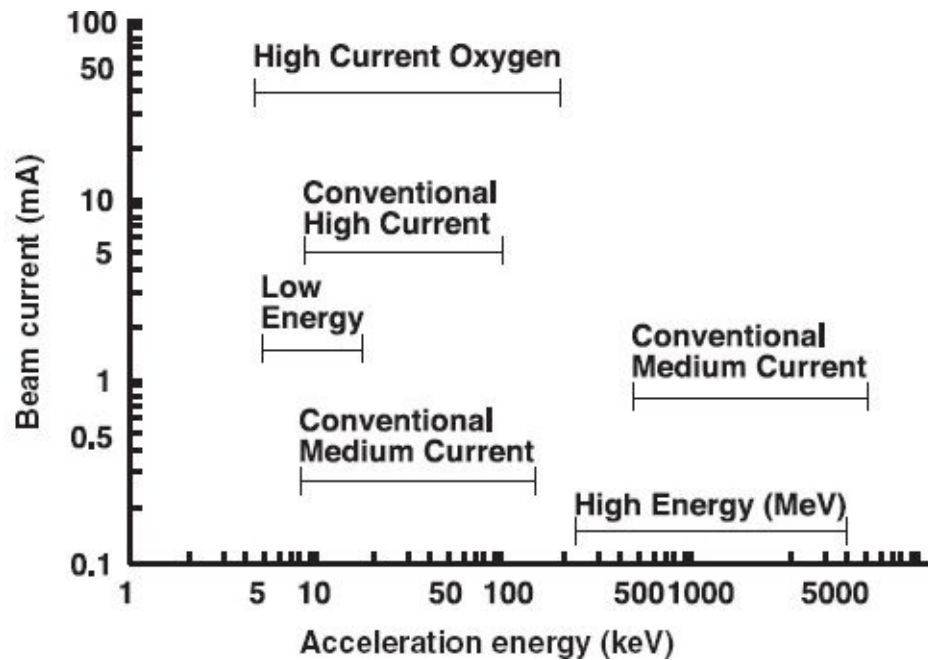


FIGURE 11.24 Ion implanters.

The conventional descriptions of ion implanters were based on applications. However, more advanced implanters cannot be classified easily; several systems are capable of a broader process window than the conventional descriptions indicate.

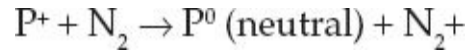
A successful ion implant relies on the implantation of only the desired dopant atom. Single-dopant implantation requires that the system be maintained at a low pressure, greater than 10^{-6} torr. The danger is that any residue molecules in the system (such as air) can become accelerated and end up in the wafer surface. Either oil-diffusion or cryogenic high-vacuum pumps are employed to reduce the pressure. The operation of these systems is described in [Chap. 12](#).

Beam Focus

On exiting the acceleration tube, the beam separates due to repulsion of like charges. The separation (or defocus) causes uneven ion density and nonuniform layers in the wafer. For successful implantation, the beam must be focused. Electrostatic or magnetic lenses are used to focus the ions into a small diameter beam or a band of parallel beams.¹³ Parallel beams are extremely important, especially for transistor-gate applications because beam deviations can cause uneven dopant doses affecting the transistor performance.

Neutral Beam Trap

Despite the vacuum removal of the majority of the air in the system, there are still some residual gas molecules in the vicinity of the ion beam. Collisions between the ions and the residual gas atoms result in a neutralization of the dopant ion.



In the wafer, these “neutrals” cause nonuniform doping and, because they cannot be “counted” by evaluation equipment, they result in incorrect counting of the amount of dopants in the wafer. Suppression of the neutral beam is accomplished by bending the ion beam (Fig. 11.25) with electrostatic plates, leaving the neutral beam to travel straight ahead away from the wafers.

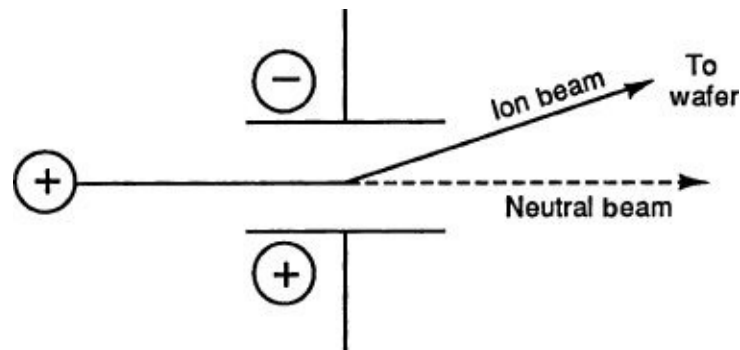


FIGURE 11.25 Deflection of ion beam from neutral beam.

Beam Scanning

Ion beams have a smaller diameter than the wafer (~1 cm). Covering the entire wafer with a uniform doping requires scanning the beam across the wafer. Three methods are used: beam scanning, mechanical scanning, and shuttering, alone or in combination.

A beam-scanning system (Fig. 11.26) has the beam pass between a number of electrostatic plates. Negative and positive charges can be controllably changed on the plates to attract and repel the ionized beam. By manipulating the charges in two dimensions, the beam can be swept across the entire wafer surface in a raster scan pattern.

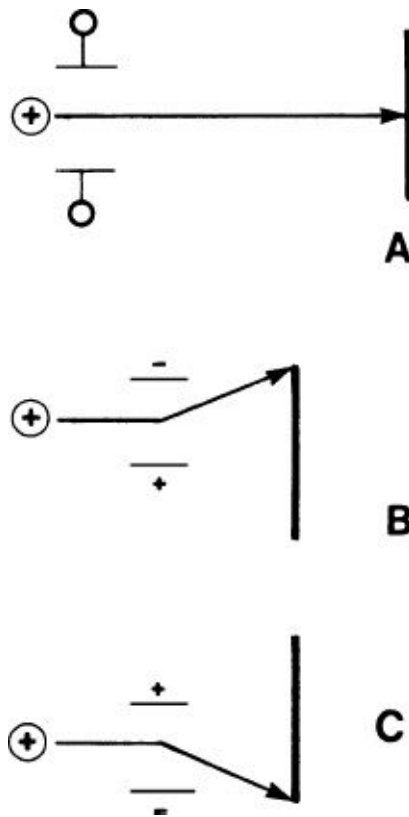


FIGURE 11.26 Electrostatic beam scanning.

Beam sweeping is used primarily in medium-current machines for single-wafer implanting. The procedure is fast and uniform. A drawback is the requirement that the beam be moved completely off the wafer to make the turns. For a large-diameter wafer, this procedure can lengthen the implant time by 30 percent or more. Another problem occurs with high-current machines where the high density of ions causes discharges (called *space charge forces*) that destroy the electrostatic plates. Wide beams are swept across the wafer. In some systems, the wafer is rotated 90° after each sweep of the beam to ensure uniformity.¹⁴

Mechanical scanning approaches the scanning problem by holding the beam in one position and moving the wafer(s) in front of it. Mechanical scanning is used primarily on high-current machines. One advantage is that there is no wasted time to turn the beam, and the beam speed is constant. If the wafer is at an angle to the beam, nonuniform implant depths can result. But, in some cases, the wafer will be oriented at an angle to the beam. Beam shuttering employs either an electronic field or a mechanical shutter to turn the beam on when it is on the wafer and off when it is not. Most systems use a combination of beam sweeping and mechanical movement.

End Station and Target Chamber

The actual implantation takes place in the target chamber of the end station. It

includes the scanning system and the wafer load and unload mechanisms. There are several tough requirements for end stations. Wafers must be loaded into the chamber; a vacuum pulled; wafers individually placed on the holder; the implant completed; and the wafers dismantled, loaded into a cassette, and removed from the chamber.

Both batch and single-wafer designs are used to present the wafer surface to the beam ([Fig. 11.27](#)). A batch process is more efficient, but there are higher maintenance and alignment chores. For batch processing, the wafers are placed on a disk that rotates them in front of the beam as it is being scanned. There may also be a mechanism to move the disk left and right in front of the beam for additional uniformity. The multiple motions increase the dose uniformity. Single-wafer designs require more time to process a group of wafers due to the additional time to load, pull the vacuum, implant, and unload. But single-wafer systems are more easily moved to control the beam angle on the wafer. Beam angle (tilt) is required to avoid channeling down crystal planes and/or to avoid beam shadowing from the raised topography of structures already on the wafer.

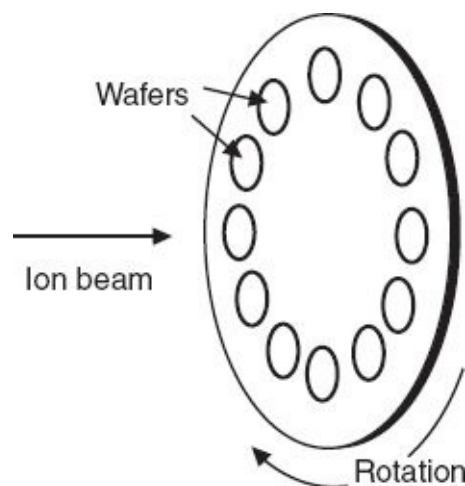


FIGURE 11.27 Mechanical scanning.

Cryogenic pumps are favored for evacuating the end station. Contamination produced during the process consists of nitrogen outgassing from the wafers and hydrogen from the photoresist masking layer. Cryogenic pumps (see [Chap. 13](#)) are a capture type and hold the potentially dangerous hydrogen frozen in the pump.

The mechanical motions can take longer than the implant itself. Improvements include load locks to allow loading without breaking the chamber vacuum. A big challenge is to keep particulate generation low during all of the mechanical movements in the chamber.¹³ Installation of antistatic devices in the chamber is critical. Electrostatic handlers (no mechanical clamps) are an option.¹⁶

Wafer breakage can cause contamination from wafer chips and dust, which in turn

requires time-consuming cleaning. Contamination on the wafers causes shadowing that blocks the ion beam. Production speed must be maintained with a system that quickly evacuates the chamber for implanting and returns it to room pressure for exit and reloading. The target chamber may house a detector (called a *Faraday cup*) to “count” the number of ions impacting the surface. These detectors can automate the process by allowing beam contact with the wafer until the correct dose is achieved.

High-current implantation can cause the wafer to heat up, and these machines often have cooling mechanisms on the wafer holders. These machines may also have an *electron flood gun* ([Fig. 11.28](#)) designed to minimize a buildup of charge on the wafer surface that can electrostatically attract contamination.

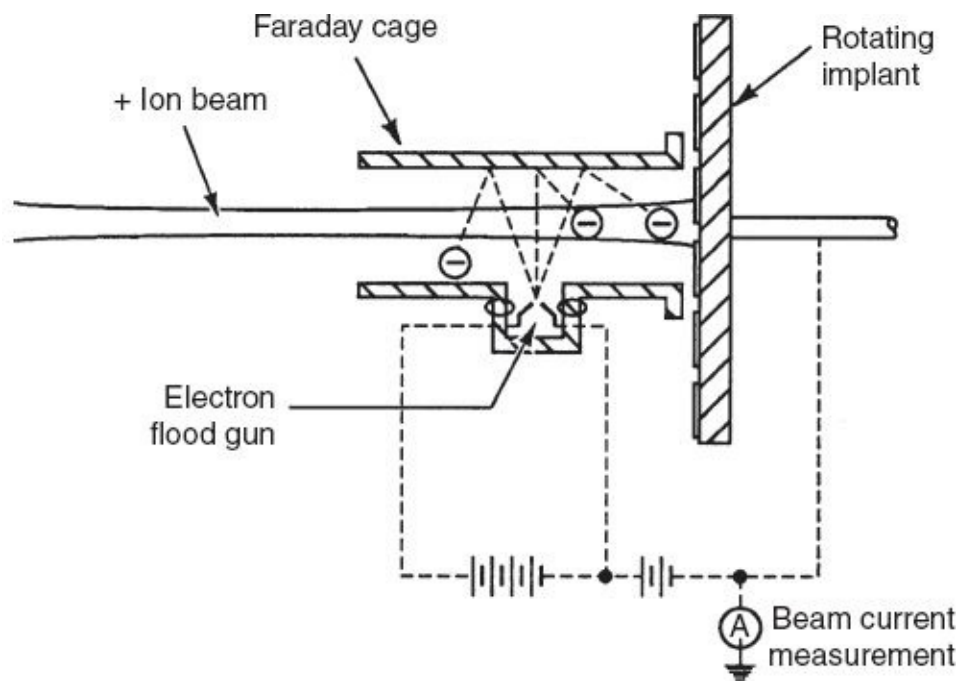


FIGURE 11.28 Electron flood gun.

Ion-Implant Masks

A major advantage of ion implantation is the variety of masks that are effective blocks to the ion beam. In diffusion doping, the only effective mask is silicon dioxide. Most films employed in the semiconductor process can be used to block the ion beam, including photoresist, silicon dioxide, silicon nitride, aluminum, and other thin metal films. [Figure 11.29](#) compares the thicknesses required to block a 200-keV implant for various dopants.

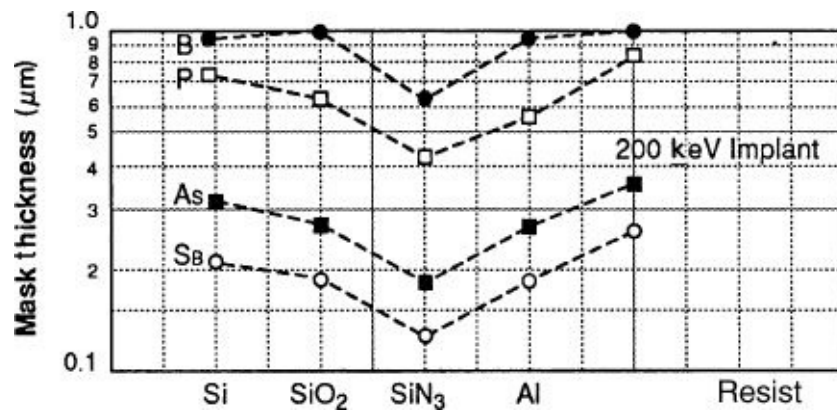


FIGURE 11.29 Barrier thickness required to block ion beam.

The use of resist films as a beam block rather than an etched opening in an oxide layer offers the same dimension control advantage as the lift-off process; the etch step and its variability are eliminated. Use of resist layers is also more productive. Options to the use of silicon dioxide increase overall yield by minimizing the number of heating steps the wafers undergo.

Dopant Concentration in Implanted Regions

The distribution of the ions in the wafer surface is different from the distribution after a diffusion process. The number and location of the dopant atoms in a diffusion process is determined by the diffusion laws and time and temperature. In an ion-implant process, the number of atoms (dose) implanted is determined by the beam current density (ions per square centimeter) and the implant. The location of the ions in the wafer is a function of the incoming energy of the ions, the orientation of the wafer, and the stopping mechanism of the ion. The first two factors are physical ones. The heavier the incoming ion and/or the higher its energy, the deeper into the wafer it will move. The wafer orientation influences the stopping position because different crystal planes have different atom densities and the incoming ions are stopped by the wafer atoms.

Within the wafer, the ions are slowed and stopped by two mechanisms. The positive ions are slowed by electronic interactions with the negatively charged electrons in the crystal. The other interaction is the physical collision with the nucleus of the wafer atoms. All of the stopping factors are variable; the ions have a distribution of energies, the crystal is not perfect, and the electronic interactions and collisions vary. The net result is that the ions come to rest over an area in the wafer (Fig. 11.30). They are centered about a depth called the *projected range* and fall off in density on each side of it. Additional implants create similar distribution patterns. Projected ranges for different dopants are shown in Fig. 11.31. The mathematical shape of the ion distribution is a Gaussian curve. A junction between the implanted

ions and the bulk doping in the wafer takes place where the ion concentration equals the bulk doping concentration.

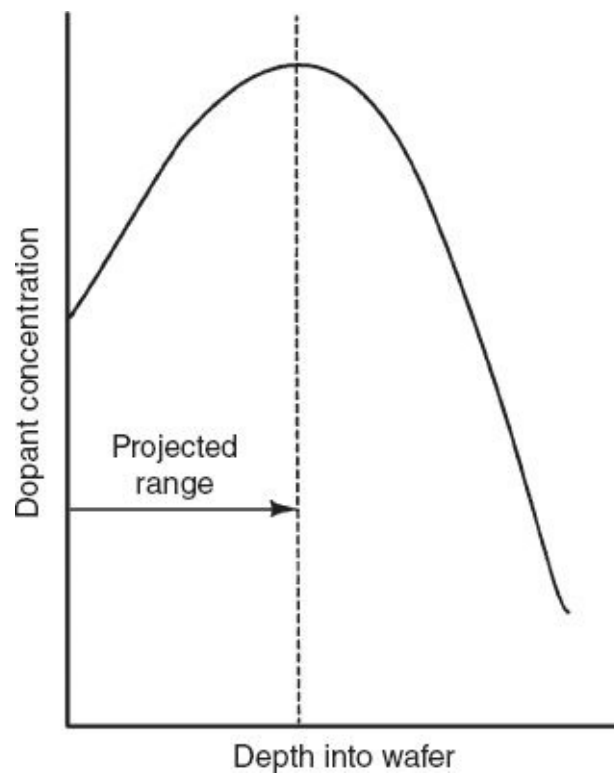


FIGURE 11.30 Dopant concentration profile after an ion implant.

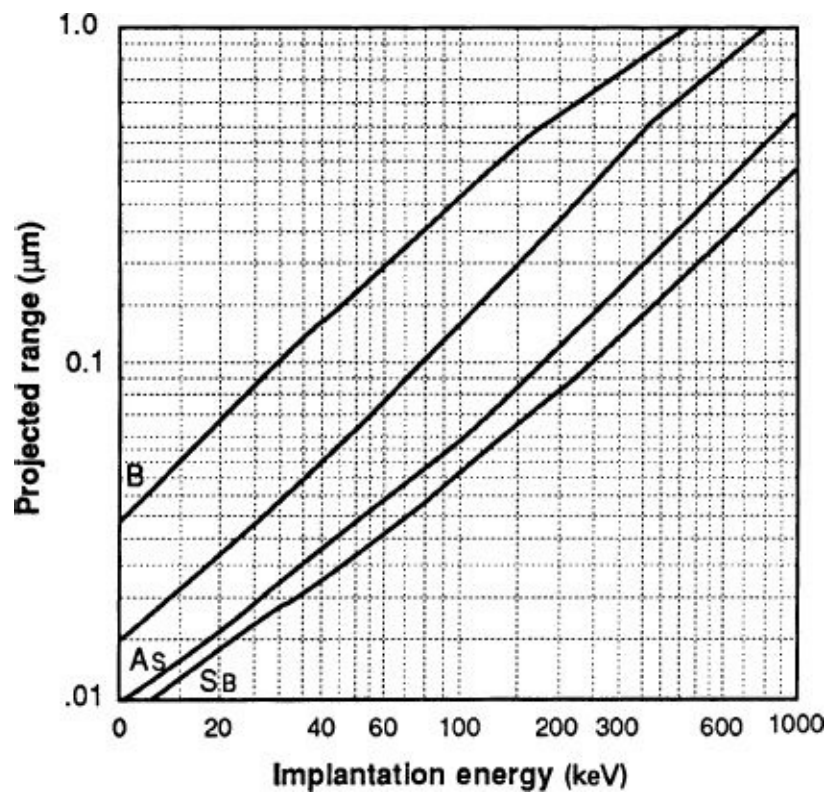


FIGURE 11.31 Projected range of various dopants in silicon. (After Blanchard, Trapp, and Shepard.)

Crystal Damage

During the process of implantation, the wafer crystal structure is damaged by the colliding ions. There are three types of damage: lattice damage, damage cluster, and vacancy interstitial.¹⁷ Lattice damage occurs when the ions collide with host atoms and displace them from their lattice site. A damage cluster occurs when displaced atoms in turn displace other substrate atoms, creating a cluster of displaced atoms. The most common implant-produced defect is a vacancy-interstitial. This defect comes about when an incoming ion knocks a substrate atom from a lattice site and the displaced atom comes to rest in a nonlattice position (Fig. 11.32).

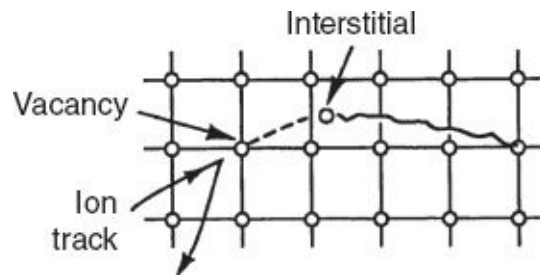


FIGURE 11.32 Vacancy-interstitial damage mechanism.

Light atoms, such as boron, produce a small number of displaced atoms. The heavier atoms, phosphorus and arsenic, generate a large number of displaced atoms. With prolonged bombardment, the regions of dense disorder may change to an amorphous (noncrystal) structure. In addition to the structural damage to the wafer from ion implantation, there is an electrical effect. The electrical characteristics in the damaged regions are because the implanted atoms do not occupy lattice sites.

Annealing and Dopant Activation

Restoration of the crystal damage and electrical activation of the dopants can be achieved by a thermal heating step. The temperature of the anneal is below the diffusion temperature of the dopant to prevent lateral diffusion. A typical anneal in a tube furnace will take place between 600 and 1000°C in a hydrogen atmosphere.

RTP techniques are also used for postimplant annealing. RTP offers fast surface heating that restores the damage without the substrate temperature rising to the level where atoms move due to diffusion effects. Additionally, the anneal can take place in seconds, whereas a tube process takes 15 to 30 min.

Channeling

The crystalline structure of the wafer presents a problem during the ion-implantation process. The problem comes about when the major axis of the crystal wafer are parallel to the ion beam. Ions can travel down the channels, reaching a depth as much as 10 times the calculated depth. An ion-concentration profile of a channeled cross-section ([Fig. 11.33](#)) shows an significant amount of additional dopants. Channeling is minimized by several techniques: a blocking amorphous surface layer, misorientation of the wafer, and creating a damage layer in the wafer surface.

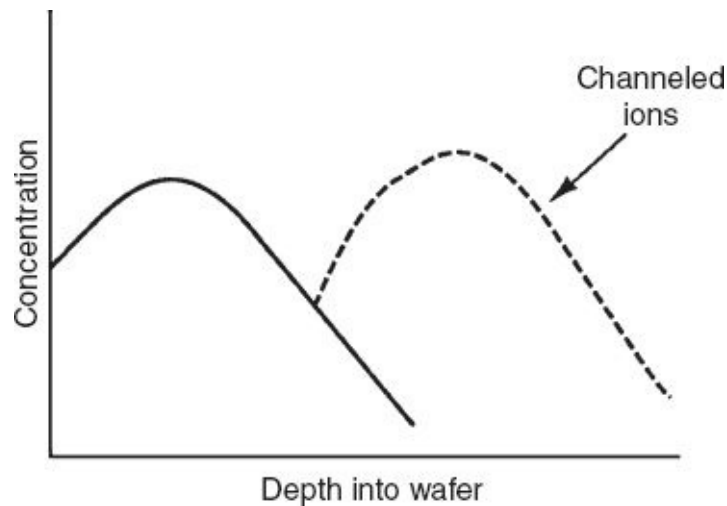


FIGURE 11.33 Effect of channeled ions on total dose.

The usual blocking amorphous layer is simply a thin layer of grown silicon dioxide ([Fig. 11.34](#)). The layer randomizes the direction of the ion beam so that the ions enter the wafer at different angles and not directly down the crystal channels. Misorientation of the wafer 3 to 7° off the major plane also has the effect of preventing the ions from entering the channels ([Fig. 11.35](#)). Predamaging the wafer surface with a heavy silicon or germanium implant creates a randomizing layer in the wafer surface ([Fig. 11.36](#)). The method increases the use of the expensive ion-implant tool. Channeling is more of a problem with low-energy implants and heavy ions.¹⁸

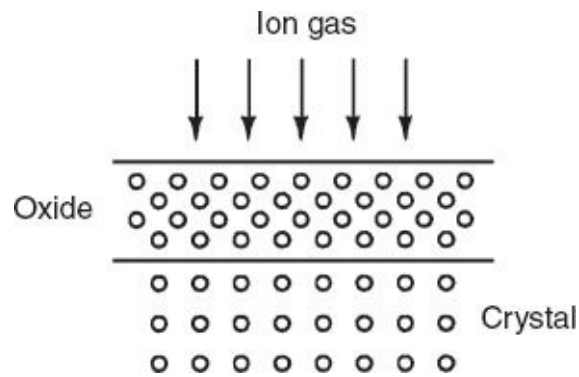


FIGURE 11.34 Implant through an amorphous oxide layer.

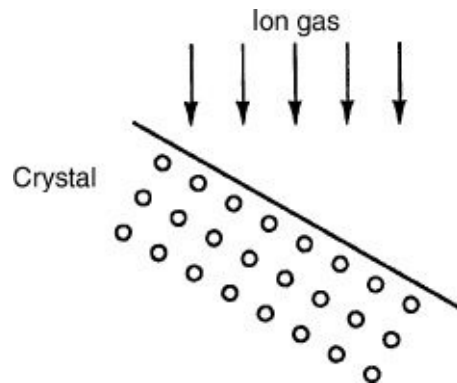


FIGURE 11.35 Misorient the beam direction to all crystal axes.

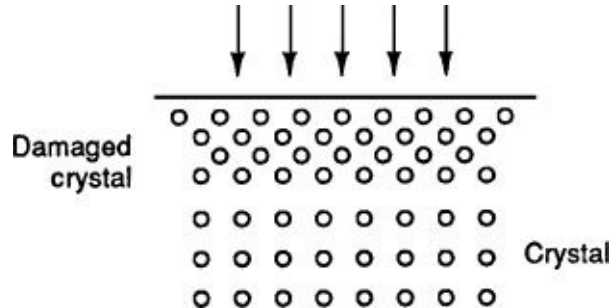


FIGURE 11.36 Predamage on the crystal surface.

Evaluation of Implanted Layers

The evaluation of implanted wafers is essentially the same as for diffused layers. Four-point probes are used to determine the sheet resistance of the layer. Spreading resistance techniques and capacitance-voltage techniques determine the concentration profile, dose, and junction depth. Junction depths can also be determined by bevel and decoration methods. These procedures are explained in [Chap. 14](#).

For implanted layers, a special structure, called the *van der Pauw structure*, is sometimes used in place of a four-point probe ([Fig. 11.37](#)). The structure allows a determination of sheet resistance without the contact resistance problems of a four-point probe. Variation in the implanted wafers can come from many sources: the

beam uniformity, variations in voltage, scanning variations, and problems in the mechanical systems. These potential problems have the possibility of causing a wider sheet resistance surface variation than a diffusion process. To detect and control the sheet resistance across the wafer surface, mapping techniques are popular and, for critical implants, required. Wafer surface maps ([Fig. 11.38](#)) are drawn from four-point probe measurements that are computer corrected for proximity and edge effects.

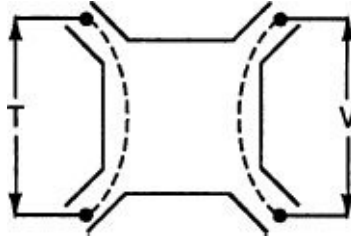


FIGURE 11.37 Van der Pauw test pattern.

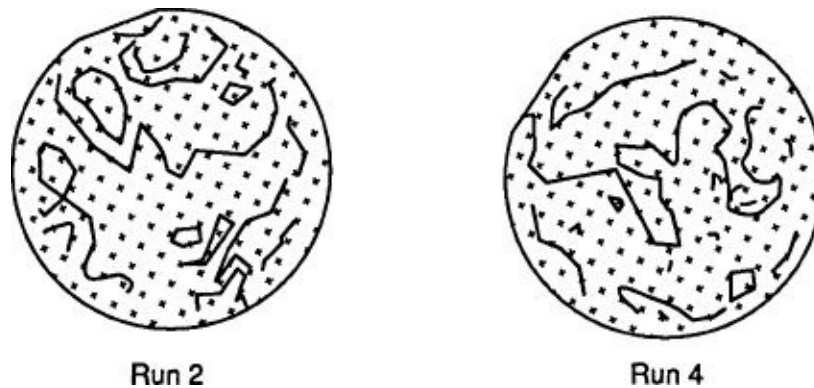


FIGURE 11.38 Four-point probe surface measurement pattern. (Courtesy of Prometrics.)

A measurement technique unique to ion implantation is optical dosimetry. The technique requires spinning a glass disk with a layer of photoresist. Before being placed in the ion implanter, the resist film is scanned with a dosimeter that measures the absorption of the film. The information is stored in a computer. The wafer and film receive the same implant as the device wafers. The resist absorbs the ion dose and darkens. After the implant, the film is again scanned. The computer subtracts the before-implant value for each location and prints out a contour map of the surface. The spacing of the contour lines is an indication of the uniformity of the dopants in the surface ([Fig. 11.39](#)).

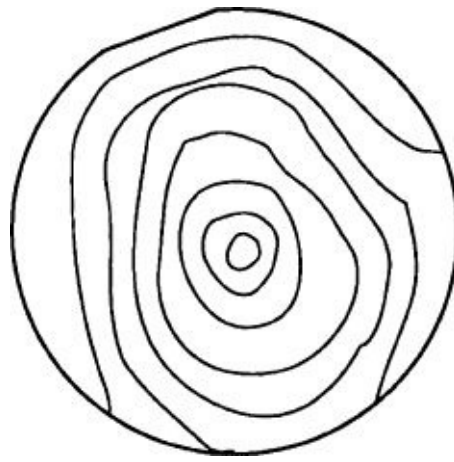


FIGURE 11.39 Contour map of ion dose.

Uses of Ion Implantation

An ion implantation can be substituted for any diffusion-deposition process. The greater control and lack of side diffusion make it the preferred doping technique for dense and small-feature-size circuits. A predeposition use in CMOS devices is the creation of the deep P-type wells (see [Chap. 16](#)), called *retrograde wells*, using high-energy implants.

A particular challenge is ultra-shallow junctions. These are junctions in the sub-125-nm range. As devices are continually scaled to smaller dimensions, junction areas also become reduced. This in turn leads to lower energy ion-implant processes to minimize surface damage and channeling. This leads to implanting pure boron instead of BF_3 , which contains corrosive fluorine. All of these requirements are being met in a new generation of implanters that can deliver high-dose beams at acceptably low energies.

One of the most important uses of ion implantation is for MOS gate threshold adjustment ([Fig. 11.40](#)). A MOS transistor consists of three parts: a source, a drain, and a gate. During operation, a voltage is applied between the source and drain regions. However, no current can flow between the regions until the gate becomes conductive. The gate becomes conductive when a voltage applied to it causes a conductive channel to form in the surface and connects the source and drain. The amount of voltage required to first form the connecting channel is called the *threshold voltage* of the device. The threshold voltage is very sensitive to the dopant concentration in the wafer surface under the gate. Ion implantation is used to create the required dopant concentration in the gate region. Also, in MOS technology, ion implantation is used to alter the field dopant concentration. However, in this use, the purpose is to set a concentration level that prevents current flow between adjacent devices. In this application, the implanted layer is part of a device *isolation* scheme.

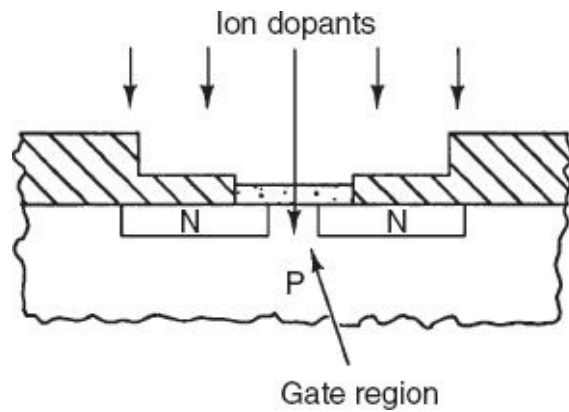


FIGURE 11.40 Ion doping of MOS gate region.

In bipolar technology, ion implantation is used to create all of the various transistor parts. The “custom” dopant profiles available from implantation can improve device performance. One particular application is arsenic-buried layers. When the buried layer is diffused, the high concentration of arsenic atoms affects the quality of a subsequent epitaxial layer deposited on the surface. By using an arsenic ion implant, a high concentration of arsenic is possible, and an anneal process restores the damage, allowing a higher-quality epitaxial layer to be deposited.

Resistors for both MOS and bipolar circuits are good candidates for ion implantation. Diffused resistors vary in uniformity from 5 to 10 percent, whereas ion-implanted resistors have variations of less than 1 percent or better. [Figure 11.41](#) is a table of typical applications of ion-implantation doping.

Typical Uses of Ion Implantation
• Gate threshold voltage adjustment • Ultra-shallow junctions • Buried layers (retrograde wells) • Predeposition layer • Insulation layer of SOI

FIGURE 11.41 Ion-implantation uses.

The Future of Doping

Ion implantation also has its drawbacks. The equipment is expensive and complex. Training and maintenance take longer and have more stringent requirements. The machines present new dangers in the form of high voltages and more toxic gas use. From the process perspective, the biggest worry is over the ability of the annealing processes to completely eliminate the implant-induced damage. However, despite

the drawbacks, ion implantation is the preferred doping process for advanced circuits.¹⁹ In addition, many new structures are possible only with the unique advantages offered by ion implantation. Another technique is *plasma ion immersion* (PII). In this technique the analyzing magnet is removed from the system. Dopants leave the source section and a plasma field increases their energy.²⁰ This technique places the wafer in a plasma field (similar to ion milling or sputtering) in the presence of dopant atoms. With proper charges on the dopant atoms and the wafer, the dopant atoms accelerate to the wafer surface and penetrate much like an ion implantation. The difference is the lower energy of the plasma field yielding less wafer charging and giving more control of the shallow junctions.²¹

Ion implantation, in one form or another, is the doping technology that is carrying the industry into the nanometer era. The benefits include: • Precise dose control from 10^{10} to 10^{16} atom/cm²

- Uniform topping of large areas
- Dopant profile control through energy selection
- Relative ease of implanting all dopant elements
- Minimal side diffusion
- Implantation of nondopant atoms
- Capability of doping through surface layers
- Choice of dopant barriers for selective doping
- Tailored doping profiles in deep (retrograde) wells

Review Topics

Upon the completion of this chapter, you should be able to:

1. Define an N-P junction.
2. Draw a flow diagram of a complete diffusion process.
3. List the three most common dopants used in silicon technology.
4. List the three types of deposition sources.
5. Draw a typical concentration-versus-distance curve for a deposition and drive-in.
6. List the major parts of an ion implanter.
7. Describe the principle of an ion implanter.
8. Compare the advantages and disadvantages of diffusion and ion-implant processes.

References

1. Griffin, P. B., and Plummer, J. D., "Advanced Diffusion Models for VLSI," *Solid State Technology*, May 1988:171.
2. Robinson, K. T., "A Guide to Impurity Doping," *Micromanufacturing and Test*, April 1986:52.
3. Guise, P., and Blanchard, R., *Modern Semiconductor Fabrication*, 1986, Reston Books, Reston, VA:46.
4. Felch, S., "A Comparison of Three Techniques for Profiling Ultrashallow p⁺-n Junctions," *Solid State Technology*, PennWell Publishing Company, Jan. 1993:45.
5. Saraswat, K., *EE 311/Shallow Junctions*, <http://www.stanford.edu/class/ee311/NOTES/ShallowJunctions>, June 2013.
6. Rubin, L., and Poate, J., *Ion Implantation in Silicon Technology*, American Institute of Physics, www.aip.org/tip/INPHFA/vol-9/iss-3/p12.html, June 2013.
7. Amem, M., Berry, I., Class, W., et al., *Ion Implantation, Handbook of Semiconductor Manufacturing Technology*, 2007, CRC Press, Hoboken, NJ:7–46.
8. Burggraaf, P., "Ion Implanters: Major Trends," *Semiconductor International*, Apr. 1986:78.
9. Iscoff, R., "Are Ion Implanters the Newest Clean Machines?" *Semiconductor International*, Cahners Publishing, Oct. 1994:65.
10. Cheung, N., "Ion Implantation," *Semiconductor International*, Cahners Publishing, Jan. 1993:35.
11. England, J., "Charge Neutralization during High-Current Ion Implantation," *Solid State Technology*, PennWell Publishing, July 1994:115.
12. Japan Report, *Semiconductor International*, Cahners Publishing, Nov. 1994:32.
13. Eaton Corp., Product Video, The NV8200P, 1993.
14. Ibid.
15. Iscoff, R., "Are Ion Implanters the Newest Clean Machines?" *Semiconductor International*, Cahners Publishing, Oct. 1994:65.
16. "Wafer Handler for Ion Implanters, Varian Semiconductor Equipment," *Solid State Technology*, PennWell Publishing, Jul. 1994:131.
17. Hayes, J., and Van Zant, P. *Doping Today Seminar Manual*, Semiconductor Services, 1985, San Jose, CA.
18. Zrudsky, D., "Channeling Control in Ion Implantation," *Solid State Technology*, Jul. 1988:73.
19. Cheung, N., "Ion Implantation," *Semiconductor International*, Cahners Publishing, Jan. 1993:35.

[20](#). Braun, A., "Ion Implantation Goes Beyond Traditional Parameters," *Semiconductor International*, Mar. 2002:48.

[21](#). Singer, P., "Plasma Doping: An Implant Alternative?" *Semiconductor International*, Cahners Publishing, May 1994:34.

CHAPTER 12

Layer Deposition

Introduction

Doped regions and N-P junctions are the electronic hearts of the active components in a semiconductor transistor. However, it takes various other layers of semiconductors, dielectrics, and conductors to complete the components and facilitate the integration of the components into the circuit. These layers are added to the wafer surface by a number of techniques. The principle ones are chemical vapor deposition (CVD), physical vapor deposition (PVD), electroplating, spin-on, and evaporation. This chapter describes the most commonly used CVD techniques and dielectric and semiconductor materials deposited on the wafer surface. PVD, electroplating, spin-on, and evaporation processes are described in [Chap. 13](#).

Advances in photomasking technology have allowed the fabrication of ever-smaller dimensioned ultra large-scale integration (ULSI) circuits. But as the circuits have shrunk, they also have grown in the vertical direction through increased numbers of added layers. In the 1960s, bipolar devices had two layers deposited by CVD, an epitaxial layer, and a top-side passivation layer of silicon dioxide ([Fig. 12.1](#)), while early metal oxide semiconductor (MOS) devices had just a passivation layer ([Fig. 12.2](#)). By the 1990s, advanced devices featured four levels of metal interconnects, requiring numerous deposited layers. The “stack” has grown with more metal layers and device schemes, and insulating layers. Specific metallization techniques are addressed in [Chap. 13](#) and common device structures described in [Chap. 17](#).

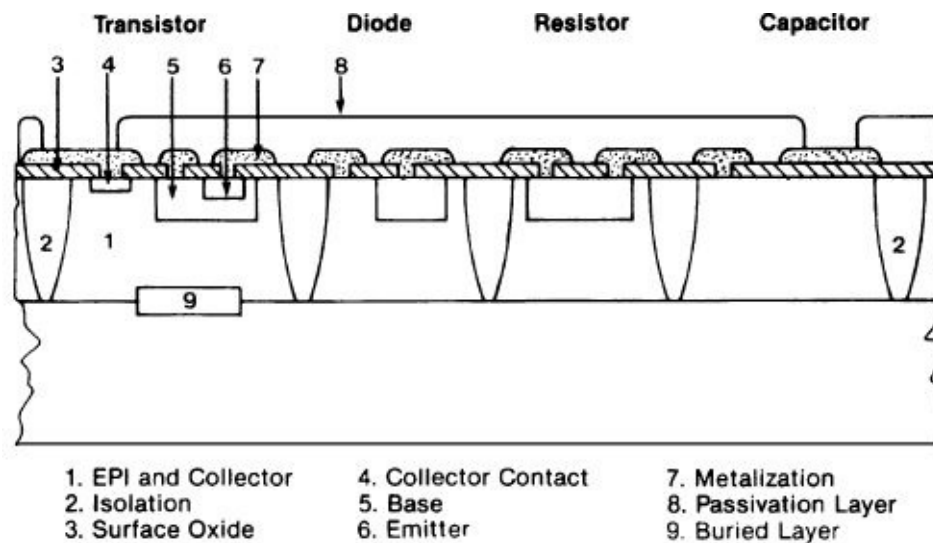


FIGURE 12.1 Cross-section of bipolar circuit showing the epitaxial layer and isolation.

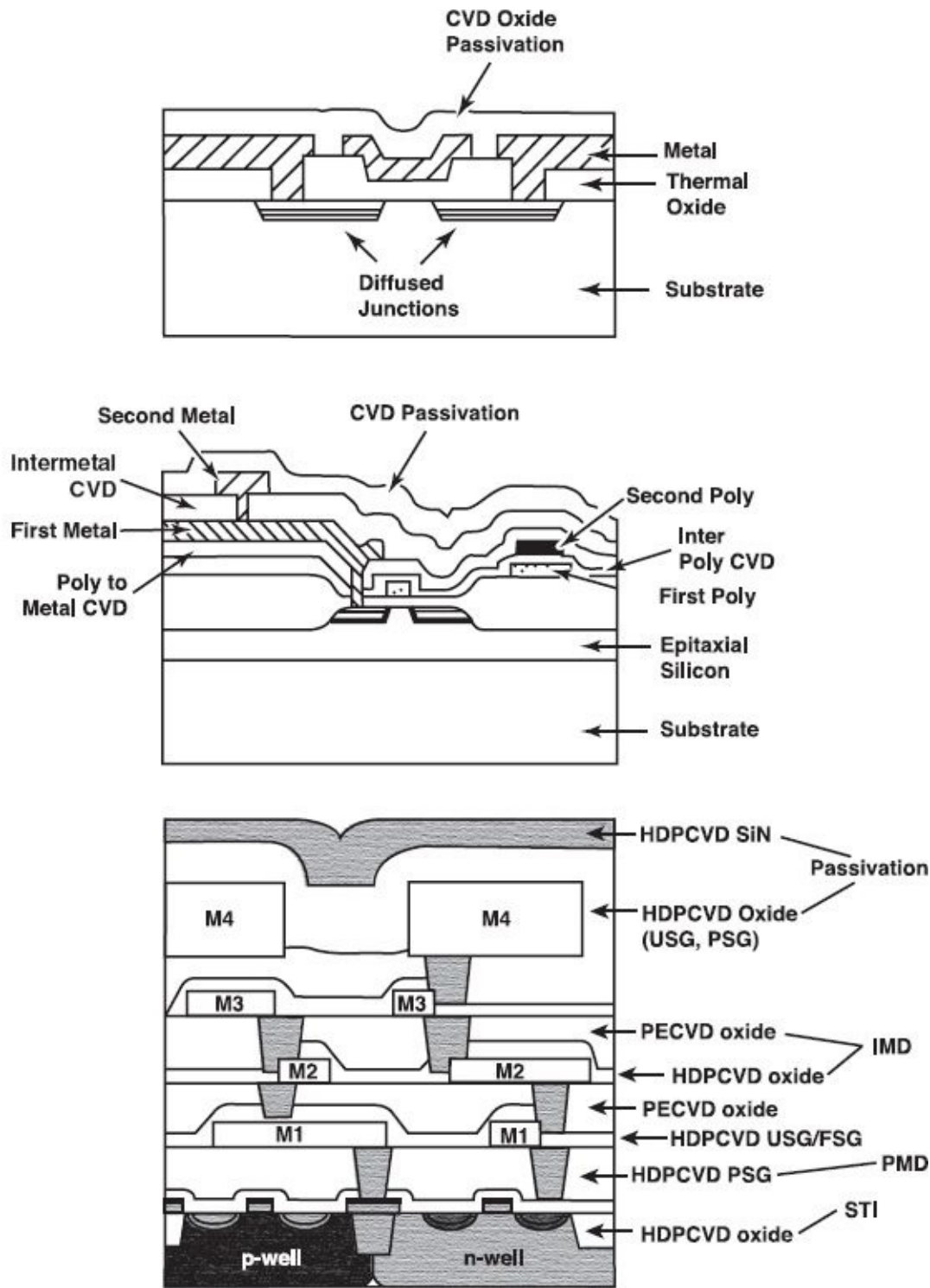


FIGURE 12.2 Evolution of MOS layers.

The added layers take a variety of roles in the device or circuit structures. They include the following:

- Deposited doped silicon layers called *epitaxial layers* (see related section in this chapter)
- Intermetal dielectrics (IMDs)

- Vertical (trench) capacitors
- Intermetal conducting plugs

- Metal conducting layers
- Final passivation layers

There are two primary techniques for layer deposition: chemical vapor deposition and physical vapor deposition. The metallization deposition techniques of sputtering and evaporation are explained in [Chap. 13](#) as is the dual-damascene process. The uses of the particular films, while indicated in this chapter, are detailed in [Chaps. 16](#) and [17](#). Chemical vapor deposition, the subject of this chapter, is practiced in a number of atmospheric and low-pressure techniques.

Film Parameters

Device layers must meet general and specific parameters. The specific parameters are noted in the sections on individual layer materials. General criteria that all films must meet for semiconductor use include:

- Thickness or uniformity

- Surface flatness or roughness
- Composition or grain size
- Stress-free
- Purity
- Integrity

Uniform thickness is required of films to meet both electrical and mechanical specifications. Deposited films must be continuous and free of pinholes to prevent the passage of contamination and to prevent shorting of sandwiched layers. This is of great importance for thin films. Recall that the thickness of a layer is one of the factors contributing to its resistance. Also, thinner layers tend to have more pinholes and less mechanical strength. Of particular concern is the maintenance of thickness oversteps ([Fig. 12.3](#)). Excessive thinning at a step can cause electrical shorts and/or unwanted induced charges in the device. The problem becomes very acute in deep and narrow holes and trenches, called *high-aspect-ratio patterns*. The ratio is calculated by dividing the depth by the width ([Fig. 12.3](#)). One problem is a thinning of the deposited film at the lip of the trench. Another is thinning in the bottom of the trench. Filling high-aspect trenches is a major issue in the execution of multimetall structures.

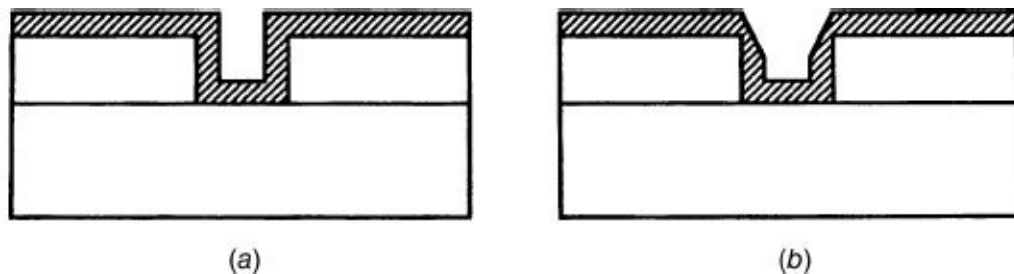


FIGURE 12.3 Thinning of deposited layer (b) at step.

Surface flatness is as important as the thickness. In [Chap. 10](#), the effect of steps and surface roughness on image formation was detailed. Deposited films must be flat and as smooth as the material and deposition method will allow to minimize steps, cracking, and subsurface reflections.

Deposited films must be of the desired uniform composition. Some of the

reactions are complex, and it is possible that films will be deposited with something other than the intended *stoichiometry*. Stoichiometry is the methodology by which the quantities of reactants and products in chemical reactions are determined. In addition to chemical composition, grain size is important. During deposition, the film materials tend to collect or grow into grains. Varying grain size within films of the same composition and thickness will yield different electrical and mechanical properties. This is because electrical current flow is affected as it passes through the grain interfaces. Mechanical properties also change with the size of the grain interface area.

Stress-free films are another requirement. A film deposited with excess stress will relieve itself by forming cracks. Cracked films cause surface roughness and can allow contamination to pass through to the wafer. In the extreme, they cause electrical shorts.

Purity (i.e., no unwanted chemical elements or molecules in the film) is required for the film to carry out its intended function. For example, oxygen contamination of an epitaxial film will change its electrical properties. Purity also includes the exclusion of mobile ionic contaminants and particulates.

An electrical parameter of importance to deposited films is capacitance (see [Chap. 2](#)). Semiconductor metal conduction systems need high conductivity and, therefore, low-resistance and low-capacitance materials. These are referred to as *low-k dielectrics*. Dielectric layers used as insulators between conducting layers need high capacitances or *high-k dielectrics*.

Chemical Vapor Deposition Basics

Not surprisingly, the growth in the number and kinds of deposited films has resulted in a number of deposition techniques. Where the process engineer of the 1960s had a choice of only atmospheric CVD, today's engineer has many more options ([Fig. 12.4](#)). These techniques are described in the following sections.

Atmospheric Pressure (AP)	Low Pressure (LP) and Ultra High Vacuum (UHV)
Cold wall • Horizontal • Vertical • Barrel • Vapor phase epitaxy Metalorganic CVD	Hot wall - Plasma enhanced - Vertical isothermal - Molecular beam epitaxy (MBE)

FIGURE 12.4 Overview of deposition systems.

Thus far, the terms *deposition* and *CVD* have been used without explanation. In semiconductor processing, deposition refers to any process in which a material is physically deposited on the wafer surface. Grown films are those, such as silicon dioxide, that formed from the material in the wafer surface. The majority of films are deposited by a CVD technique. In concept, the process is simple ([Fig. 12.5](#)). Chemicals (C) containing the atoms or molecules required in the final film are mixed and reacted in a deposition chamber to form a vapor (V). The atoms or molecules deposit (D) on the wafer surface and build up to form a film. [Figure 12.6](#) illustrates the reaction of silicon tetrachloride (SiCl_4) with hydrogen to form a deposited layer of silicon on the wafer. Generally, CVD reactions require energy to take place.

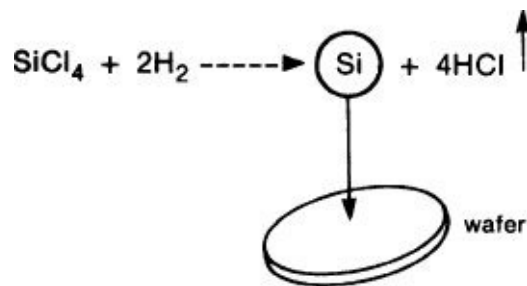


FIGURE 12.5 Chemical vapor deposition of silicon from silicon tetrachloride.

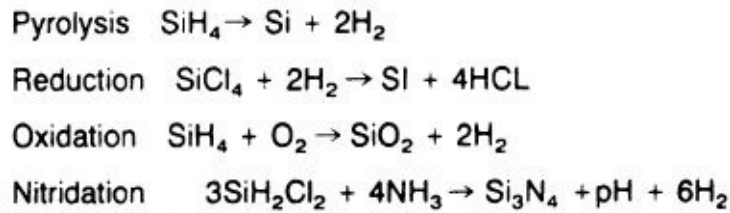


FIGURE 12.6 Examples of CVD reactions.

The chemical reactions that take place fall into the four categories of pyrolysis, reduction, oxidation, and nitridation ([Fig. 12.6](#)). *Pyrolysis* is the process of chemical reaction driven by heat alone. *Reduction* causes a chemical reaction by reacting a molecule with hydrogen. *Oxidation* is the chemical reaction of an atom or molecule with oxygen. *Nitridation* is the chemical process of forming silicon nitride by exposing a silicon wafer to nitrogen at a high temperature.

Deposited films grow in several distinct stages ([Fig. 12.7](#)). The first stage, *nucleation*, is very important and critically dependent on substrate quality. Nucleation occurs as the first few atoms or molecules deposit on the surface. These first atoms or molecules form islands that grow into larger islands. In the third stage, the islands spread, finally coalescing into a continuous film. This is the transition stage of the film growth with a typical thickness of several hundred angstroms. The transition region film has chemical and physical properties much different from those of the final, thicker “bulk” film.¹

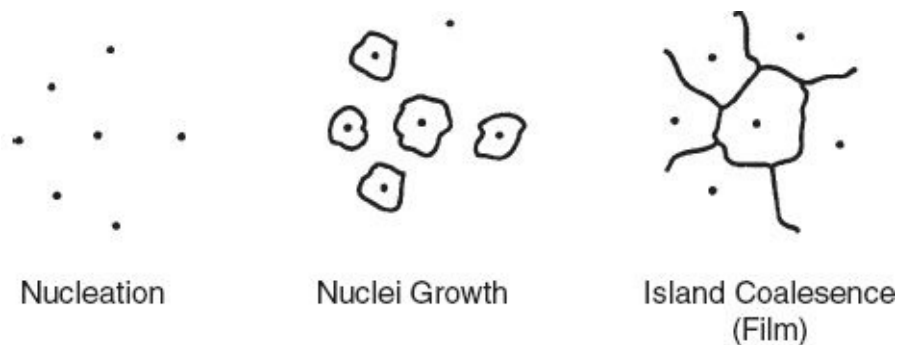


FIGURE 12.7 CVD film growth steps.

After the transition film is formed, the bulk growth begins. Processes are designed to produce three different structures: amorphous, polycrystalline, and single crystal ([Fig. 12.8](#)). These terms have been defined previously. A poorly defined or controlled process can result in a film with the wrong structure. For example, attempting to grow a single-crystal epitaxial film on a wafer with islands of unremoved oxide will result in regions of polysilicon in the bulk film.

Amorphous



Polycrystalline



Single
Crystal

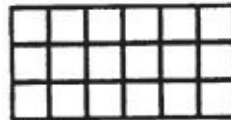


FIGURE 12.8 Types of film structure.

Basic CVD System Components

CVD systems come in a wide variety of designs and options. Understanding the many variations is helped by an examination of the basic subsystems common to most CVD systems ([Fig. 12.9](#)). In most respects, a CVD system has the same basic parts as a tube furnace (described in [Chap. 7](#)): source cabinet, reaction chamber, energy source, wafer holder (boat), and loading and unloading mechanisms. In some cases, the CVD system is a tube furnace identical to those used for oxidation and diffusion. The source chemicals are housed in a source section. Vapors are generated from pressurized gas cylinders or liquid source bubblers. Gas flow control is maintained by pressure regulators, mass-flowmeters, and timers.²

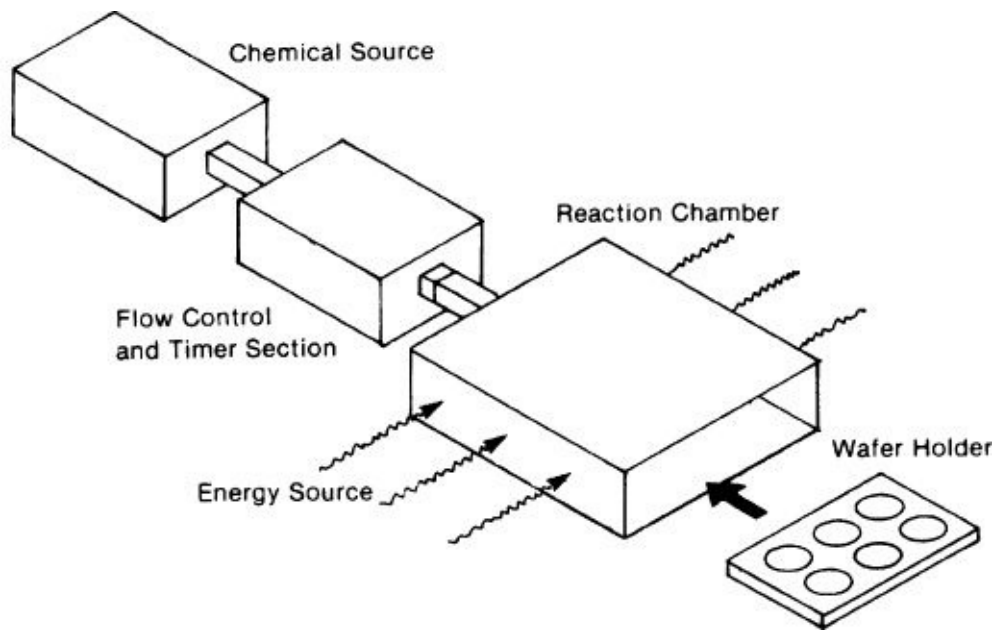


FIGURE 12.9 Basic CVD subsystems.

The actual deposition takes place on the wafers in a reaction chamber. Energy sources can be heat conduction, convection, induction RF, radiant, plasma, or ultraviolet. Energy sources are explained in the sections on particular systems. Temperatures range from room temperature to 1250°C, depending on the reaction, film thickness required, and the growth parameters.

The fourth basic part of the system is the wafer holder. Different chamber configurations and heat sources dictate the style and material of the holders. Most production-level systems for ULSI circuits are automated from wafer load to unload. A full production system will have an associated cleaning section or station and a loading area.

CVD Process Steps

A CVD process follows the same steps as an oxidation: preclean (and etch, if required), deposition, and evaluation. Cleaning processes are those already described to remove particulates and mobile ionic contaminants. Chemical vapor deposition, like an oxidation, takes place in cycles. First, the wafers are loaded into the chamber, usually with an inert atmosphere. Next, the wafers are brought to temperature. Chemical vapors are introduced for as long as required to deposit the film. Finally, the chemical source vapors are flushed out and the wafers removed. Evaluation of the films is done for thickness, step coverage, purity, cleanliness, and composition. Evaluation techniques are explained in [Chap. 14](#).

CVD System Types

CVD systems ([Fig. 12.10](#)) are divided into two primary types: atmospheric-pressure (AP) and low-pressure (LP). There are a number of atmospheric pressure CVD systems (APCVD). However, most films for advanced circuits are deposited in systems where the pressure has been lowered. These are called low-pressure CVD or LPCVD.

Atmospheric systems	Low-pressure systems
Hot wall	Hot Wall
Cold wall	Plasma enhanced
• Epitaxial	• Hot-wall tube
• Barrel	• Cold-wall planar—batch
• Pancake	• Cold-wall planar—single
• Single wafer	
	High-density plasma
	Atomic layer deposition

FIGURE 12.10 CVD reactor designs.

Another differentiation is cold wall versus hot wall. Cold-wall systems directly heat the wafer holder or wafers, with induction or radiant heating. The walls of the chamber remain cold (or cooler). Hot-wall systems heat the wafers, the wafer holder, and the chamber walls. The advantage of cold-wall CVD is that the reaction occurs only at the heated wafer holder. In a hot-wall system, the reaction occurs throughout the chamber, leaving reaction products on the inside chamber walls. The reaction products build up, necessitating rigorous and frequent cleaning to avoid contaminating the wafers.

CVD systems are operated with two principal energy sources: thermal and plasma. Thermal sources are tube furnaces, hot plates, and RF induction. Plasma-enhanced chemical vapor deposition (PECVD) in combination with lower pressure offers the unique advantage of lowered temperatures and good film composition and coverage.

A specialty CVD used to deposit compound films, such as GaAs, is vapor phase epitaxy (VPE). A newer technique used to deposit metals is a metalorganic CVD (MOCVD) source in a VPE system. The last deposition method described is the non-

CVD molecular beam epitaxy (MBE) used for low-temperature deposition of thin films in a very controlled process.

Atmospheric-Pressure CVD Systems

As the name implies, atmospheric CVD system reactions and deposition take place at atmospheric pressure. There are a number of system designs that fall under this category ([Fig. 12.10](#)).

Horizontal-Tube Induction-Heated APCVD

The first widespread use of CVD was for the deposition of silicon epitaxial films for bipolar devices. The basic system design is still used ([Fig. 12.11](#)). It is essentially a horizontal tube furnace, but with some significant differences. First, the tube has a square cross-section. The major difference, however, is in the heating method and wafer holder.

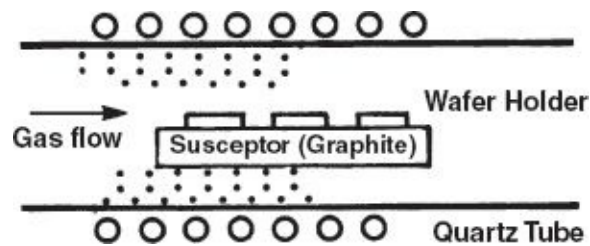


FIGURE 12.11 Cold-wall induction APCVD with horizontal susceptor.

Wafers are arranged on a flat graphite slab and positioned in the tube. Surrounding the tube are copper coils that are connected to an RF generator. The RF waves traveling in the coils pass through the quartz tube and the flowing gas in the tube without heating them. This is the cold-wall aspect of the system. When the radiant waves reach the graphite wafer holder, they “couple” with the molecules of the holder, causing the graphite to heat up. This heating method is called *induction*.

The heat of the holder is passed to the wafers by conduction. The film deposition takes place at the wafer surface (and at the holder surface). One problem with this type of system is downstream depletion of the reactants in the laminar gas flow. Laminar flow is required to minimize turbulence. But if the wafers are laid flat in the chamber, the layer of gas closest to the wafers becomes depleted. This results in successively thinner films along the wafer holder. A wafer holder tilted in the tube corrects the problem ([Fig. 12.12](#)).

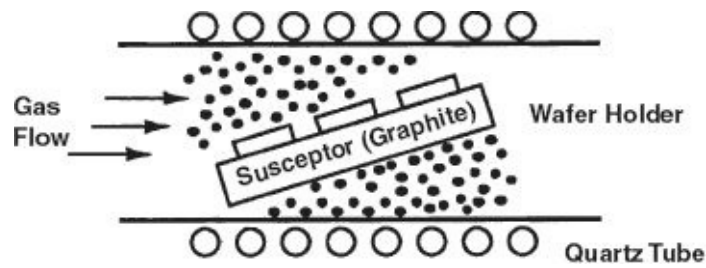


FIGURE 12.12 Cold-wall induction APCVD with tilted susceptor.

Barrel Radiant-Induction-Heated APCVD

Larger-diameter wafers laid horizontally on a holder in a horizontal system have a low packing density. And larger wafer holders strain the uniformity capabilities of the system. The development of the barrel radiant-heated system ([Fig. 12.13](#)) solved these problems. The reaction chamber of the system is a cylindrical stainless steel barrel with high-intensity quartz heaters placed about the inside surface. The wafers are placed on a graphite holder that rotates in the center of the barrel. The rotation of the wafers produces a more uniform film thickness as compared to horizontal systems.

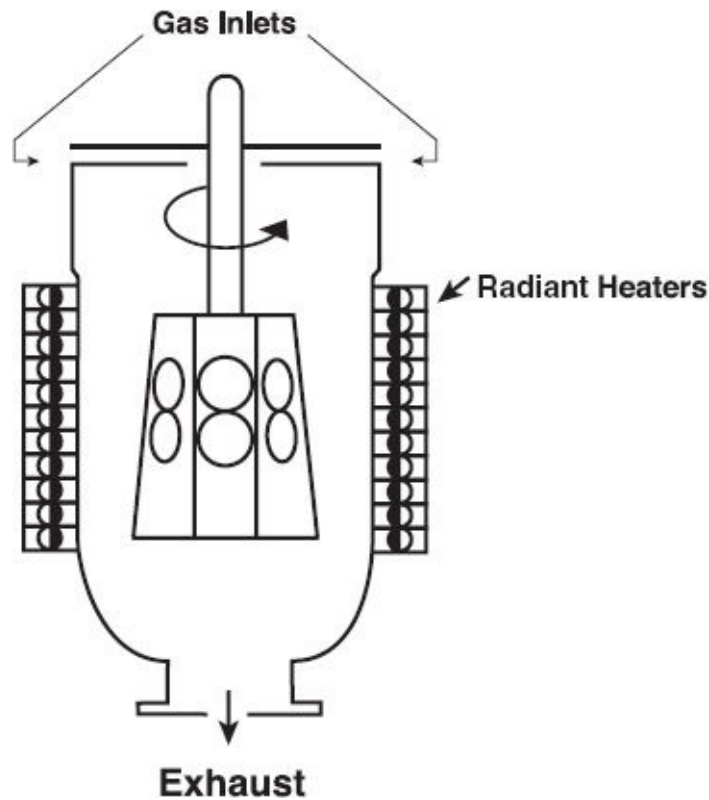


FIGURE 12.13 Cylindrical or barrel system.

Radiant heat from the lamps heats the wafer surface, where the deposition takes place. While some heating of the chamber walls occurs, the system is close to a

cold-wall deposition. Direct radiant heating produces a very controlled and even film growth. In an induction-heated system, the wafers are heated from the bottom, and, as the film grows, there is some small but measurable drop in temperature at the film surface. In the barrel system, the wafer surface is always facing the lamps and receives a more uniform temperature and film growth rate.

In 1987, Applied Material introduced a jumbo barrel system for large-diameter wafers featuring an induction heating system.³ A principal advantage of the barrel reactor is an increased productivity based on the increased number of wafers per cycle. This system configuration is popular for deposition of epitaxial silicon in the 900 to 1250°C range.

Pancake Induction-Heated APCVD

The pancake or vertical-flow APCVD system has been a favorite for small fabrication lines and R&D labs ([Fig. 12.14](#)). The wafers are held on a rotating holder of graphite and heated by induction-conduction from an RF coil below the holder. The reaction gases are fed through a tube exiting above the wafers. Vertical gas flow offers the advantage of a continuous supply of fresh reactants to the wafers, thus minimizing downstream depletion. The combination of the rotation and vertical flow of the gases produces good film uniformity. Productivity in smaller systems is restricted, as in the horizontal tube system, by the number of wafers that the pancake can accommodate.

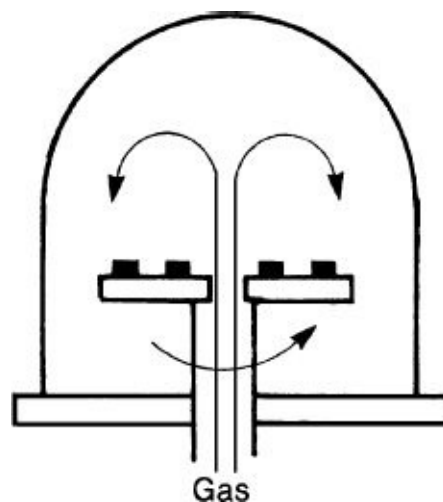


FIGURE 12.14 Rotating pancake APCVD.

A production-level variation of the pancake design is produced by Gemini Research. The reactor features radiant-resistance heating and a large-capacity holder with robot autoloading.⁴

Continuous Conduction-Heated APCVD

A horizontal conduction-heated APCVD system features mixing the gases outside the chamber and showering them onto the wafers. In one design, a heated hot-plate wafer holder moves back and forth (Fig. 12.15) under a series of nozzles that dispense a vapor of the desired material. Another version of the system (Fig. 12.16) has wafers moving on a belt under a plenum that dispenses the reacted gases.

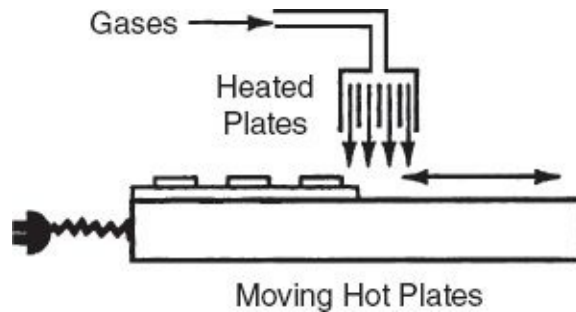


FIGURE 12.15 Moving hot-plate APCVD.

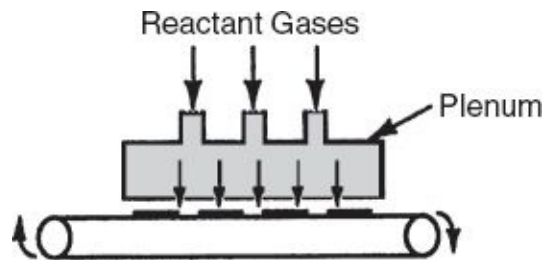
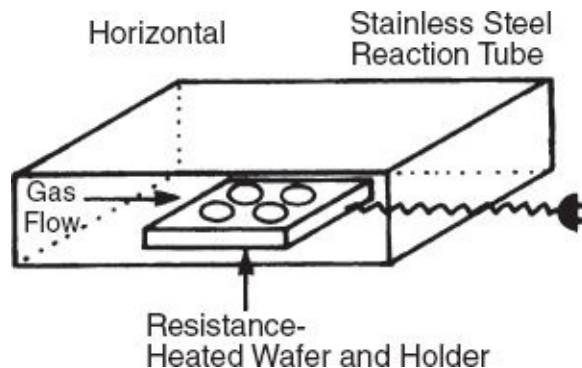


FIGURE 12.16 Continuous hot-plate APCVD.

Horizontal Conduction-Heated APCVD

One of the original CVD designs is the horizontal conduction-heated APCVD system (Fig. 12.17) used to deposit the silicon dioxide passivation layers. Wafers are mounted on a removable hot plate and placed in a stainless steel chamber. The hot plate heats the wafers and chamber walls (hot-wall system). Reaction gases are fed into the chamber.



Low-Pressure Chemical Vapor Deposition

Uniformity and process control within atmospheric-pressure CVD systems rely on temperature control and the flow dynamics in the system. A factor influencing film uniformity and step coverage is the mean free path of the molecules in the reaction chamber. The mean free path is the average distance a molecule will travel before colliding with another object in the chamber, be it another molecule, the wafers, or wafer holder. Collisions change the direction of the particles. The longer the mean free path, the higher the uniformity of the film deposition. A major determiner of the length of the mean free path is the pressure in the system. Lowering the pressure in the chamber increases the mean free path and the film uniformity. Decreasing the pressure also allows a lowering of the deposition temperature.

These benefits became available to the industry in 1974 when Unicorp, under license to Motorola Inc., introduced an LPCVD system that operated at a few hundred millitorr.⁵ The list of LPCVD system advantages includes:

- Lower chemical reaction temperature

- Good step coverage and uniformity
- Vertical loading of wafers for increased productivity and lower exposure to particles
- Less dependence on gas-flow dynamics
- Less time for gas-phase reaction particles to form
- Can be performed in standard tube furnaces

A vacuum pump must be added to the system to reduce the pressure in the chamber. A discussion of the types of pumps used with LPCVD systems appears in [Chap. 13](#).

Horizontal Conduction-Convection-Heated LPCVD

One production-level LPCVD system uses a horizontal tube furnace ([Fig. 12.18](#)), with three major exceptions. The tube is connected to a vacuum pump that pulls the system down to a pressure range of 0.25 to 2.0 torr.⁶ A second change is a ramping of the temperature in the center zone to offset reaction depletion down the tube. The third change may be special injectors at the gas inlet end to improve gas mixing and deposition uniformity. In some systems, the injectors are positioned directly over the wafers. Disadvantages of this system design are particles formed on the inside wall surface (hot-wall reactions), uniformity along the tube axis, the use of cages around the wafers to minimize particle contamination, and the higher downtime required for frequent cleaning.

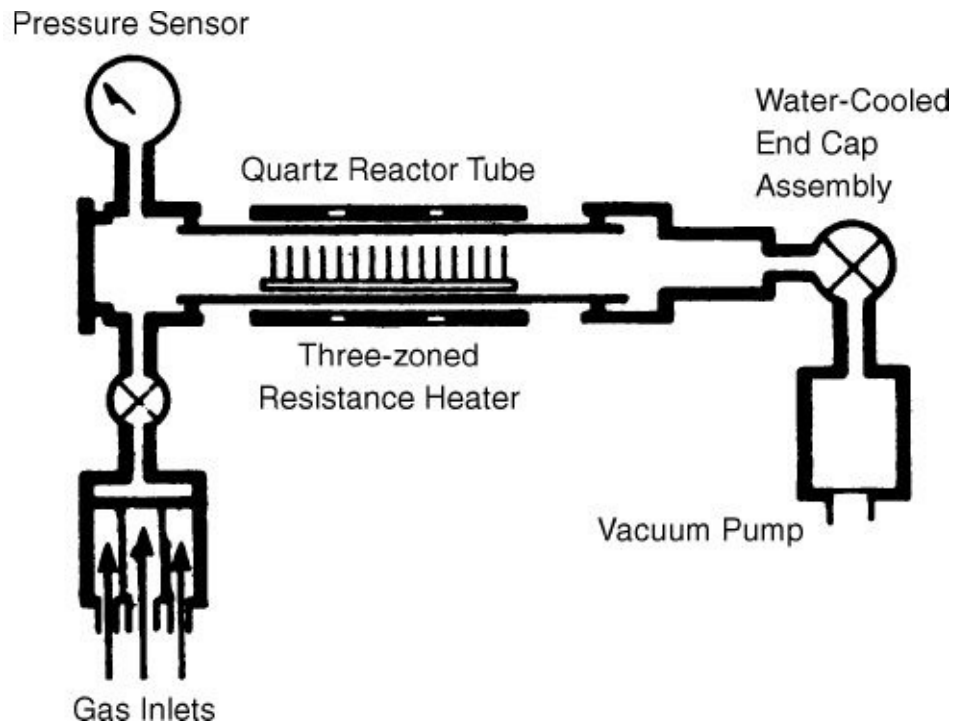


FIGURE 12.18 Horizontal hot-wall LPCVD system.

These systems are most often used for polysilicon, silicon dioxide, and silicon nitride films with typical thickness uniformities of ± 5 percent. The primary deposition variables are temperature, pressure, gas flow, gas partial pressure, and wafer spacing. These variables are carefully balanced for each deposition process. The deposition rates are somewhat lower (100 to 500 Å/min) than AP systems, but productivity is enhanced by the vertical wafer-loading densities that can approach 200 wafers per deposition.

Ultra-High Vacuum CVD

Low-temperature deposition is desirable to minimize crystal damage and lower the thermal budget that in turn minimizes the lateral diffusion of doped regions. One approach is the CVD of silicon and silicon-germanium (SiGe) in ultra-low vacuum conditions. Lowering the pressure allows keeping the deposition temperature low. Ultrahigh vacuum CVD (UHV-CVD) takes place in a tube furnace where the pressure is initially reduced to 1 to 5×10^{-9} millibar (mbar). Deposition pressure is in the 10^{-3} mbar range.²

Plasma-Enhanced CVD (PECVD)

Replacement of silicon dioxide passivation layers with silicon nitride led to the development of plasma-enhanced (PECV) techniques. A thermal silicon dioxide

deposition temperature of approximately 660°C causes unacceptable alloying of aluminum interconnects into the silicon surface (see [Chap. 13](#)). A solution to this problem was a plasma enhancement of the deposition energy. The increased energy allows a temperature under the 450°C maximum level for deposition over aluminum layers. Plasma-enhanced systems are physically similar to plasma-etch systems. They feature a parallel plate chamber operated at a low pressure. A radio-frequency-induced glow discharge or other plasma source (see [Chap. 9](#)) is used to induce a plasma field in the deposition gas. The combination of low pressure and lower temperatures provides good film uniformity and throughput.

PECVD reactors have the capability of also using the plasma for etching and cleaning the wafer prior to the deposition step. This step is the same as the dry-etch processes described in [Chap. 9](#). This *in situ* cleaning prepares the deposition surface, eliminating the problem of added contamination picked up during the loading step.

Horizontal Vertical-Flow PECVD

This system follows the design of a bottom-heated pancake vertical-flow CVD ([Fig. 12.19](#)). The plasma region is created in the top of the chamber by a radio frequency (RF) feed to an electrode or other plasma source. Wafer heating comes from radiant heaters mounted below the wafer holder, creating a cold-wall deposition system. With PECVD systems, there are several additional critical parameters to control in addition to those in a standard LPCVD reactor. They are the RF power density, the RF frequency, and the duty cycle. Film deposition speed is generally increased, but it must be controlled to prevent film stress and/or cracking.

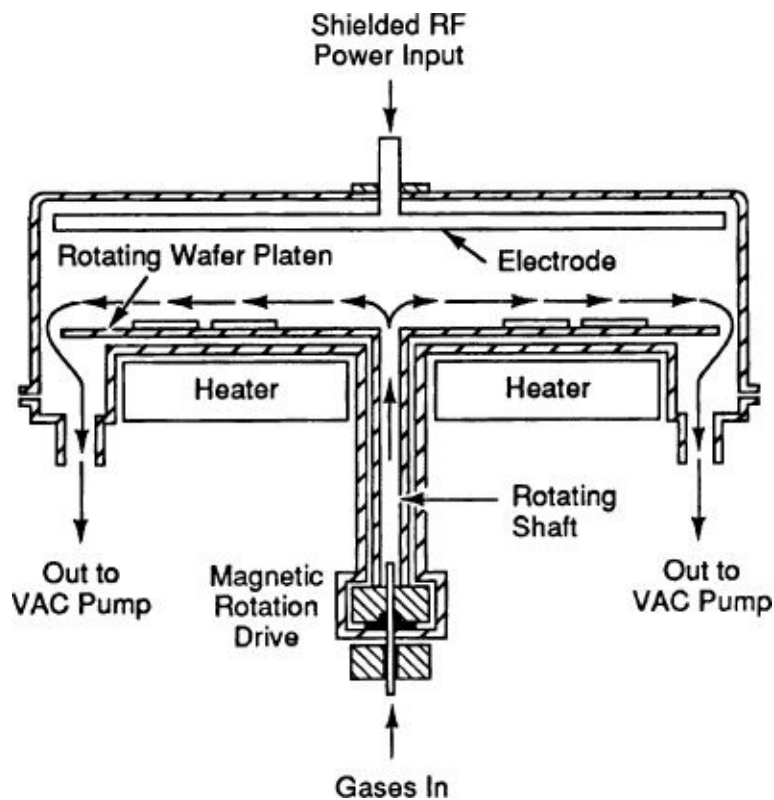


FIGURE 12.19 Vertical-flow pancake PECVD.

Another design (developed by Novellus) has the wafers sitting on a series of resistance-heated holders. The wafers are indexed around the chamber with the film built up in increments.

A single-wafer chamber PECVD system ([Fig. 12.20](#)), where the chamber is small and each successive wafer is exposed to identical conditions, addresses most of the control needs. Single-wafer systems are generally slower than batch processes. The productivity trade-off with the larger-chamber batch machines lies in fast mechanisms to feed wafers in and out of the chamber and how quickly the vacuum can be established and released. Load-lock systems enhance the productivity by moving the wafers into an antechamber, pumping it down to the required pressure, and moving the wafers into the deposition chamber as it becomes available. System designs include fixed and rotating wafer holders. Single-wafer multiple-chamber systems offer increased productivity.

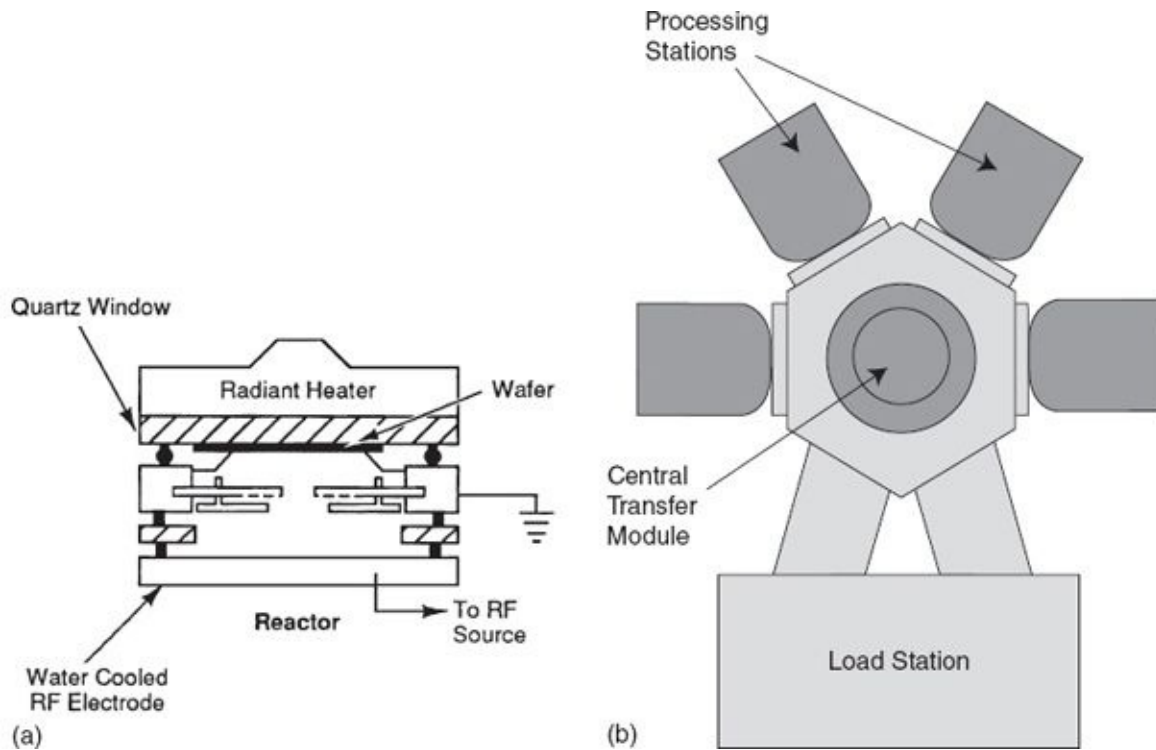


FIGURE 12.20 (a) Single-chamber planar PECVD, (b) multichamber process tool.

Barrel-Radiant-Heated PECVD

This system is a standard barrel-radiant-heated system with low-pressure and plasma capabilities. It is favored for the deposition of tungsten silicide.

High-Density Plasma CVD

Intermetal dielectric (IDL) layers are essential for multimetall structures. They also present a challenge of filling high-aspect-ratio (greater than 3:1) holes. One approach is a deposition and an *in situ* etch sequence. The first deposition exhibits the usual thinning in the bottom. Etching away the shoulder and redeposition creates a uniform layer thickness and a more planar surface.

A system to accomplish this process is *high-density plasma* CVD (HDPCVD).⁸ A plasma field is created inside a CVD chamber that contains oxygen and silane for the deposition of silicon dioxide. Also included is argon that becomes energized by the plasma and is directed to the wafer surface. This is a sputtering (see “Dry Etching,” [Chap. 13](#)) action that removes material from the surface and trench. HDPCVD has the potential of depositing a variety of materials for uses as IMD layers, etch stops, and final passivation layers.

Atomic Layer Deposition

Like every other microchip process, CVD has changed as dimensions have changed. Enter *atomic layer deposition* (ALD) as the next generation of CVD systems. It is based on a basic CVD process approach but with a unique technique, pulsing. A typical CVD system introduces the precursor chemical(s) into the chamber where the desired layer material (Si, SiO₂, Si₃N₄) is deposited on the wafer surface. In ALD, the precursors are introduced into the chamber sequentially but separated by a purging gas. The effect at the surface is illustrated in [Fig. 12.21](#). Also ALD is a self-limiting process, because the reaction takes place on the wafer surface not in the chamber. Since each film stage is growing at a monolayer rate, the control is very precise. Also, the slow rate facilitates high levels of conformity to the wafer surface and dense film composition. ALD has lowered the layer thickness from the usual CVD level of 300 Å down to the 12 Å regime.⁹ The process takes place in a vacuum. A general diagrammatic system is shown in [Fig. 12.22](#).

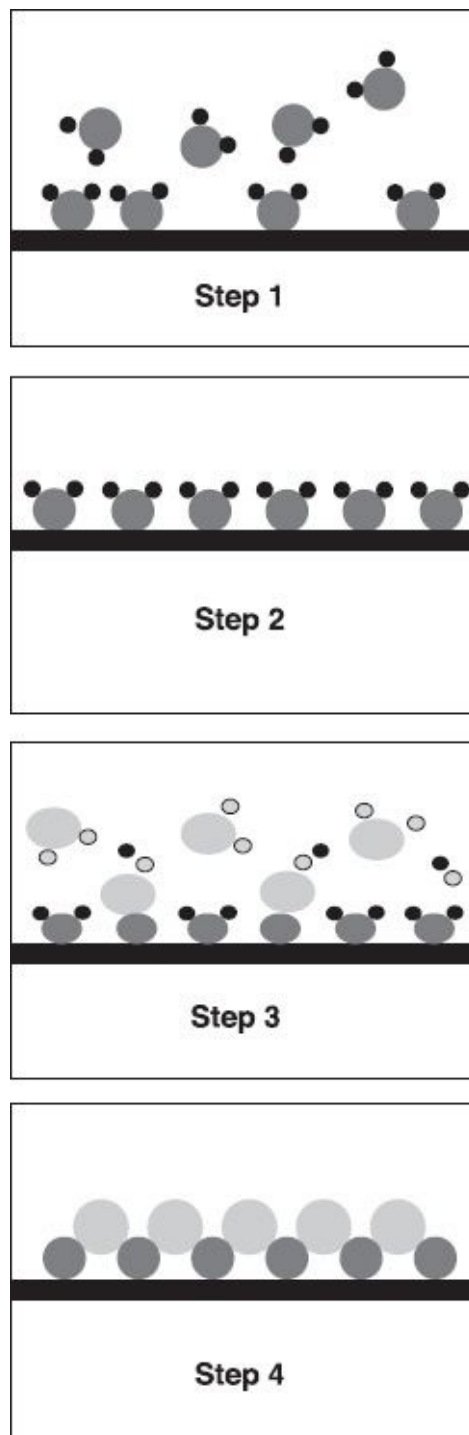


FIGURE 12.21 ALD deposition mechanism.

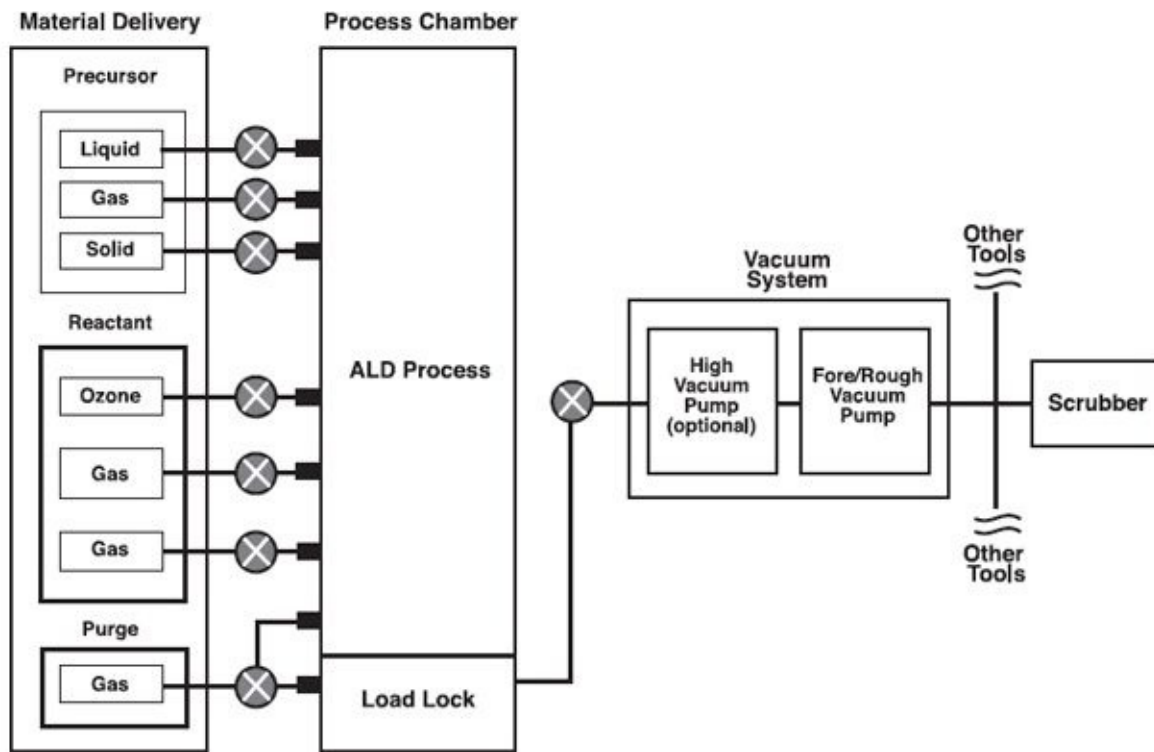


FIGURE 12.22 ALD system design.

Additionally, ALD films are very conformal and uses include very thin silicon dioxide gates, filling deep trenches with materials like aluminum oxide, and creating barrier metal layers for copper metallization processes.¹⁰ Each system offers advantages and disadvantages depending on materials and process steps. An overview is shown in [Fig. 12.23](#).

COMMON TYPES OF CVD			
Atmospheric Pressure CVD (APCVD)	Low Pressure CVD (LPCVD)	Plasma-Enhanced CVD (PECVD)	High Density Plasma CVD (HDPCVD)
APPLICATIONS			
• Low temperature oxides • Undoped silicon glass • Doped oxide in ▪ Interlayer dielectric ▪ Planarization ▪ Epitaxial layer deposition	• Barriers and etch stops • Liners - stress relief between films • High Temperature deposition of ▪ Oxides ▪ Silicon nitride ▪ Poly-Si ▪ Tungsten	• Insulators over metals • Nitride passivation • Low-k dielectrics • p-MOS gate conductor passivation • Source/drain implant stop • Premetal dielectrics • Intermetal dielectrics ▪ Gap Fills ▪ Damascene interconnected	• Shallow trench isolation filling • High aspect ration gap fill • Premetal dielectric • Intermetal dielectric ▪ Gap fills ▪ Damascene interconnected

FIGURE 12.23 Table of CVD applications.

Vapor-Phase Epitaxy

Vapor-phase epitaxy (VPE) differs from the CVD systems described in its ability to deposit compound materials, such as gallium arsenide. A VPE system^u is a combination of a standard liquid source tube furnace and a two-zone diffusion furnace. An example is the particular arrangement in [Fig. 12.24](#) used to deposit epitaxial gallium arsenide. The creation of the GaAs layer on the wafer in the main chamber proceeds in two stages. AsCl_3 is bubbled into the first section of the tube where it reacts with a solid source of gallium that is sitting in a boat. The arsenic trichloride reacts with the hydrogen in the first section to form arsenic by the reaction.

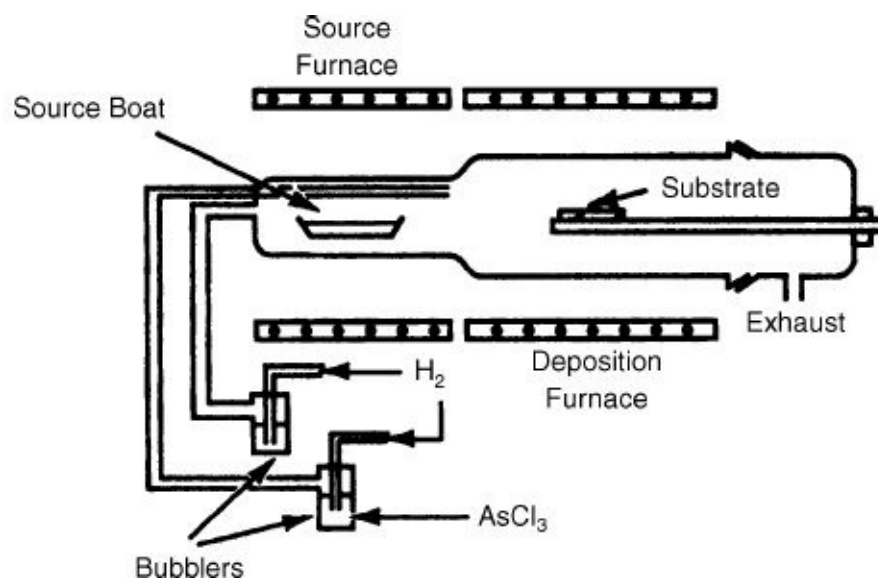
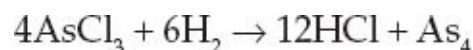
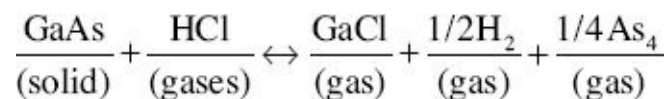


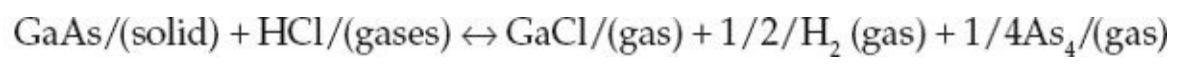
FIGURE 12.24 Diagram of gallium arsenide VPE deposition system.



The arsenic deposits on the gallium, forming a crust. The hydrogen passing over the crust reacts in the first section to form three gases that pass into the wafer section.



This section is at a somewhat lower temperature, and the reaction proceeds in reverse, depositing GaAs on the wafers. The technique offers the advantages of clean films, since the gallium and arsenic trichloride are available in very pure forms and have higher production rates than the MBE technique. On the downside, the film structures produced are not the quality of MBE films.



Molecular Beam Epitaxy

Deposition rate control, low-deposition temperature, and controlled film stoichiometry are always goals in film-deposition systems. Molecular beam epitaxy has emerged from the laboratory to claim production status as these issues have become more important. MBE is an evaporation rather than a CVD process. The system consists of a deposition chamber (Fig. 12.25) that is maintained at a low pressure to 10^{-10} torr. Within the chamber, there can be one or more cells (called *effusion cells*) that contain a very pure sample of the target material desired on the wafer. Shutters on the cells allow exposure of the wafer to the source material(s). An electron beam¹² is directed into the center of the target material, which it heats to the liquid state. In this state, atoms evaporate out of the material, exit the cell through an opening, and deposit on the wafers. If the material source is a gas, the technique is called *gas source MBE* or *GSMBE*. For most applications, the wafer in the chamber is heated to give additional energy to the arriving atoms. The additional energy fosters epitaxial growth and good film quality.

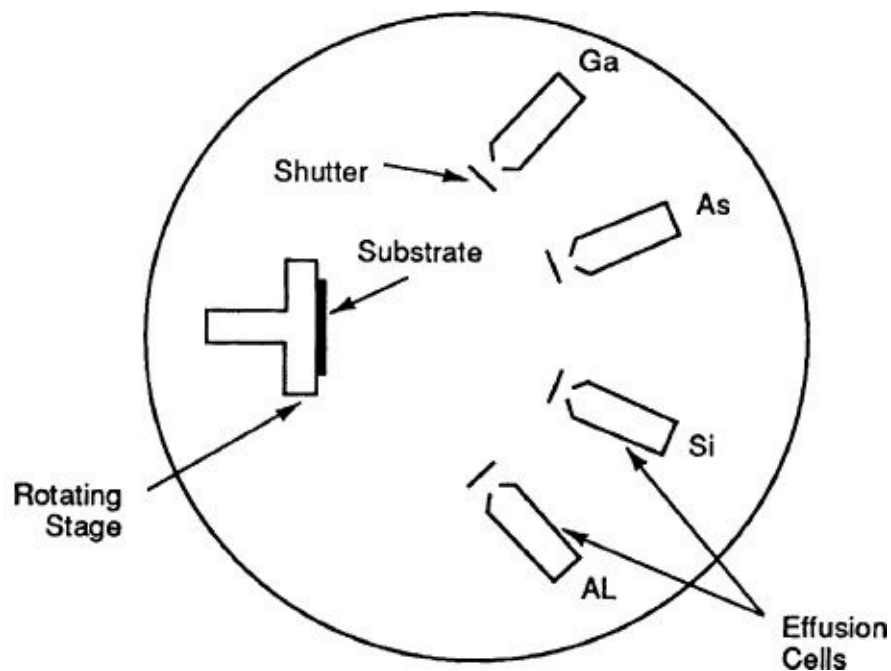


FIGURE 12.25 Diagram of an MBE deposition system.

If the wafer surface is exposed, the depositing atoms will assume the orientation of the wafer and grow an epitaxial layer. MBE offers the intriguing option of *in situ* doping by the inclusion of dopant sources in the chamber. The usual silicon dopant sources are not usable in MBE systems. Solid gallium is used for P-type doping and

antimony for N-type doping. Phosphorus deposition is virtually not possible with MBE.¹³

The primary advantage of MBE for silicon technology is the low temperature (400 to 800°C), which minimizes autodoping and out-diffusion. Perhaps the biggest advantage of MBE is the ability to form multiple layers on the wafer surface during one process step (one pump down). This option requires the mounting of several effusion cells in the chamber and shutter arrangements to direct the evaporant beams to the wafer in the right order and for the correct time.

An advantage and disadvantage of MBE is the low film growth rate of 60 to 600 Å/min.¹⁴ On the plus side, the films produced are very controllable. Films can be grown (and mixed) in one-monolayer increments. However, most semiconductor layers do not need this level of control and quality, making the low productivity and expense of the system an expensive luxury.

A bonus possibility with MBE is the incorporation of film growth and quality-analyzing instruments in the chamber. With these instruments, the process can become very controlled and produce uniform films from wafer to wafer. MBE has found production use in the fabrication of special microwave devices and for compound semiconductors such as gallium arsenide.¹⁵

Metalorganic CVD

Metalorganic CVD (MOCVD) is a newer option for the CVD of compound materials. Whereas VPE refers to a compound material deposition system, MOCVD refers to the sources used in VPE and other systems ([Fig. 12.26](#)). In a MOCVD process, desired atoms for deposition are combined with complex organic gas molecules and passed over a heated semiconductor wafer. The heated molecules break, depositing the wanted atoms on the surface, layer by layer. It can grow high-quality semiconductor layers (as thin as a millionth of a millimeter) and the crystal structure of these layers is perfectly aligned with that of the substrate.¹⁶

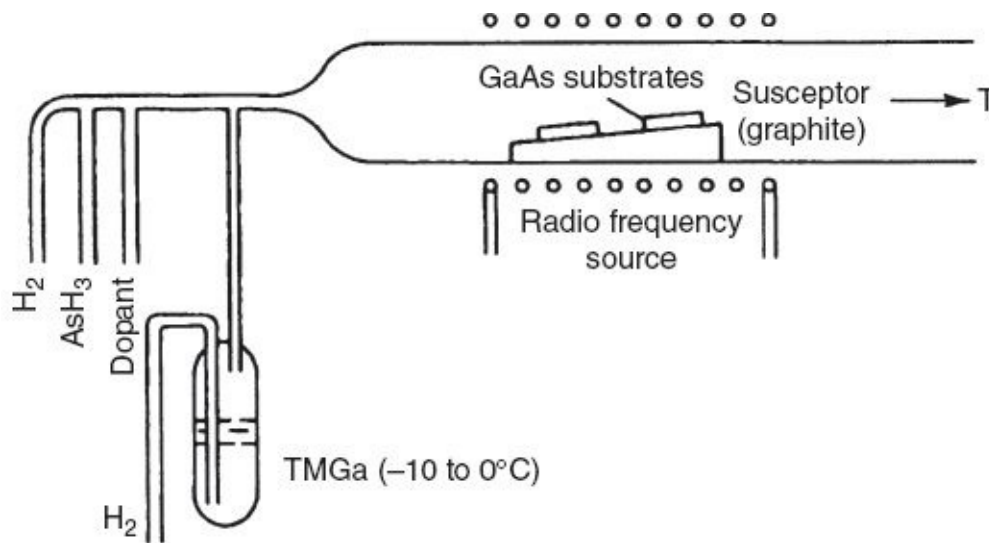


FIGURE 12.26 MOCVD system. (Source: VLSI Fabrication Principles, S. K. Ghandhi, Wiley-Interscience, 1994.)

Two chemistries are used: halides and metalorganic. The reactions described for the VPE deposition of gallium arsenide above constitute a halide process. A group III halide (gallium) is formed in the hot zone, and the III–IV compound is deposited in the cold zone. In the metalorganic process¹⁷ for gallium arsenide, trimethylgallium is metered into the reaction chamber along with arsine to form gallium arsenide by the reaction $(\text{CH}_3)_3\text{Ga} + \text{AsH}_3 \rightarrow \text{GaAs} + 3\text{CH}_4$

Where MBE processes are slow, MOCVD processes can meet volume-production requirements and accommodate larger substrates.¹⁸ MOCVD also has the capability of producing multiple layers with very abrupt changes in composition. This characteristic is critical to compound semiconductor structures. Also MOCVD, unlike MBE, can deposit phosphorus as in InGaAsP devices. Common devices made by MOCVD processes are photocathodes, high-power LEDs, long-wavelength lasers, visible lasers, and orange LEDs (see [Chap. 16](#)).

MOCVD is a broad term referring to the metalorganic chemical vapor deposit of semiconductor films. When metalorganic sources are used in a vapor-phase epitaxial system to grow epitaxial layers, the method is called MOVPE.¹⁹ Applications include metalorganic chemical-vapor deposition of epitaxial III-V semiconductor layers, including GaAs, AlAs, AlGaAs, InGaAs and InP, for both fundamental studies and device applications. Also delta doping of III-V semiconductors, growth of quantum dots, quantum wires, quantum wells, modulation-doped heterostructures, and selective-area epitaxial growth.²⁰

Deposited Films

The types of films deposited by CVD techniques are divided into their electrical classifications of semiconductors, dielectrics, and conductors. In the following sections, the deposition of the various films is examined. For each film, the principal use(s) in semiconductor devices is presented along with the particular film properties for the described use. The uses are treated in general. For a more detailed explanation of film roles in particular devices, see [Chap. 16](#). Deposition methods for conductive metal films are discussed in [Chap. 13](#).

Deposited Semiconductors

So far in this text, we have discussed the formation of wafers as the base of semiconductor devices and circuits. However, there are several drawbacks to using bulk wafers for high-quality devices and circuits. Crystal quality, doping ranges, and doping control all limit bulk wafer use. These factors placed a limit on the fabrication of high-performance bipolar transistors. A solution was found by the development of a deposited silicon layer, called an *epitaxial layer*. It is one of the major advances in the industry. Epitaxial layers were a part of the technology as early as 1950.²¹ Since then, deposited silicon layers have found additional uses in advanced bipolar device designs, a high-quality base for CMOS circuits, and silicon epitaxial layers deposited on sapphire and other substrates ([Chap. 14](#)). Gallium arsenide and other III–IV and II–VI films are also deposited in epitaxial films. Epitaxial films that are the same material as the substrate (silicon on silicon) are *homoepitaxial*. When the deposited material is different from that of the substrate (GaAs on silicon), the film is called *heteroepitaxial*.

Epitaxial Silicon

The term *epitaxial* comes from the Greek word meaning “arranged upon.” In semiconductor technology, it refers to the single-crystalline structure of the film. The structure comes about when silicon atoms are deposited on a bare silicon wafer in a CVD reactor ([Fig. 12.27](#)). When the chemical reactants are controlled and the system parameters set correctly, the depositing atoms arrive at the wafer surface with sufficient energy to move around on the surface and orient themselves to the crystal arrangement of the wafer atoms. Thus, an epitaxial film deposited on a $\langle 111 \rangle$ -oriented wafer will take on a $\langle 111 \rangle$ orientation.

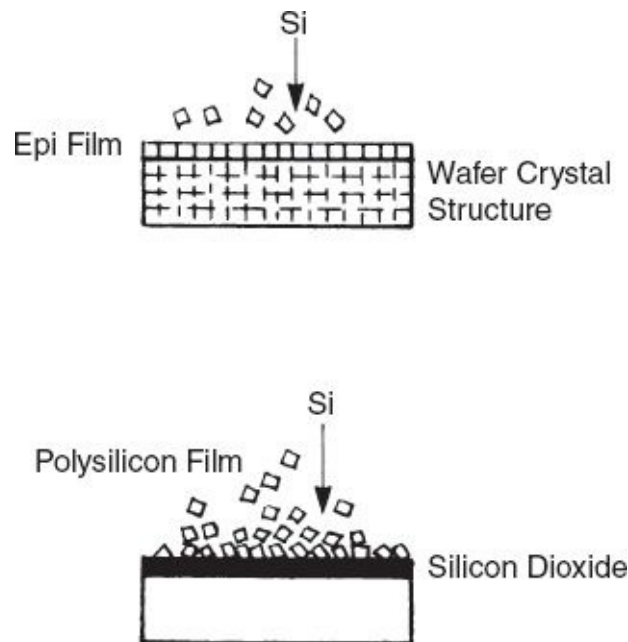


FIGURE 12.27 Epitaxial and polysilicon film growth.

If, on the other hand, the wafer surface has a thin layer of silicon dioxide, an amorphous surface layer, or contamination,²² the depositing atoms have no structure to which they can align. The resulting film structure is polysilicon. This condition is useful for some applications, such as MOS gates, and is unwanted if the goal is to grow a single-crystal film structure.

Silicon Tetrachloride Source Chemistry

A number of different sources are used for the deposition of epitaxial silicon ([Fig. 12.28](#)). Deposition temperature, film quality, growth rate, and compatibility with a particular system are factors in choosing a silicon source. An important process parameter is the deposition temperature. The higher the temperature, the faster the growth rate. Faster growth rates create more crystal defects and film cracking and stress. Higher temperatures also cause higher levels of autodoping and out-diffusion. (These effects are described in the following text.)

Silicon tetrachloride	$\text{SiCl}_4 + 2\text{H}_2 \leftrightarrow \text{Si} + 4\text{HCl}$
Silane	$\text{SiH}_4 + \text{heat} \rightarrow \text{Si} + 2\text{H}_2$
Dichlorosilane	$\text{SiH}_2\text{Cl}_2 \leftrightarrow \text{Si} + 2\text{HCl}$

FIGURE 12.28 Epitaxial silicon chemical sources.

Silicon tetrachloride (SiCl_4) is the favored source of silicon for the deposition of silicon. It allows a high-formation temperature (growth rate) and has a reversible chemical reaction. In [Fig. 12.28](#), there is a double-headed arrow, which indicates that the reaction creates silicon atoms in one direction and removes (etches) silicon in the other direction. Within the reactor, these two reactions compete with each other.

Initially, the silicon surface is etched preparing it for the deposition reaction. In the second stage, the deposition of silicon is faster than the etch, with the net result of a deposited film.

The graph in [Fig. 12.29](#) shows the effect of the two reactions. With an increasing percentage of SiCl_4 molecules in the gas stream, the deposition rate first increases. At the 0.1 ratio, the etching reaction starts to dominate and slows down the growth rate. This latter reaction is actually one of the first events in the reactor. Hydrogen chloride (HCl) gas is metered into the chamber, where it etches away a thin layer of the silicon surface, preparing it for the silicon deposition.

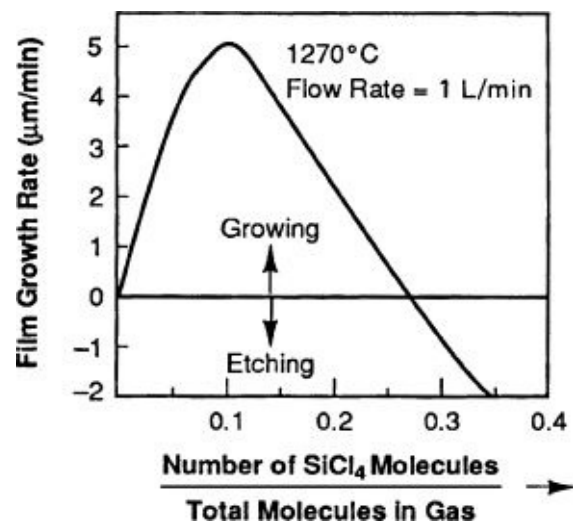


FIGURE 12.29 Growth-etch characteristics of SiCl_4 epitaxial deposition.

Silane Source Chemistry

The second-most-used silicon source chemistry is silane (SiH_4). Silane offers the advantage of not requiring a second reaction gas. It forms silicon atoms by decomposing when heated. The reaction takes place several hundred degrees lower than a silicon tetrachloride deposition, which is attractive from an autodoping and wafer-warping perspective. Also, silane does not produce pattern shift (see “Epitaxial Film Quality” below). Unfortunately, the reaction occurs at all locations in the system, creating a powdery film inside the reactor which, in turn, contaminates the wafers. Silane finds more use as a source for polysilicon and silicon dioxide depositions.

Dichlorosilane Source Chemistry

Dichlorosilane (SiH_2Cl_2) is also a lower-temperature silicon source that is used for thin epitaxial films. The lower temperature reduces autodoping and solid-state diffusion from previously diffused buried layers and provides a more uniform crystal structure.

Epitaxial Film Doping

One of the advantages of an epitaxial film is the precise doping and doping range available by the process. Silicon wafers are manufactured in a concentration range of approximately 10^{13} to 10^{19} atoms/cm³. Epitaxial films can be grown from 10^{12} to 10^{20} atoms/cm³. The upper limit is close to the solid solubility of phosphorus in silicon.

Doping in the film is achieved by the addition of a dopant gas stream to the deposition reactants. The sources of the dopant gases are exactly the same chemistries and delivery systems similar to deposition doping furnaces. In effect, the CVD deposition chamber is turned into a doping system. In the chamber, the dopants become incorporated into the growing film, where they establish the required resistivity. Both N-and P-type films can be grown on either N-or P-type wafers. The classic epitaxial film in bipolar technology is an N-type epitaxial film on a P-type wafer.

Epitaxial Film Quality

Epitaxial film quality is a prime concern of the process. In addition to the usual considerations over contamination, there are a number of faults specifically associated with epitaxial growth. Contaminated systems can cause a problem called *haze*.²³ Haze is a surface problem that varies from a microscopic disruption to severe cases that are observable as a dull matte finish. Haze comes about from residual oxygen in the reactant gases or from leaks in the system.

Contaminants on the surface at the start of the deposition result in an accelerated growth known as *spikes* (Fig. 12.30). The spikes can be as high as the film thickness. They cause holes and disruption in photoresist layers and other deposited films.

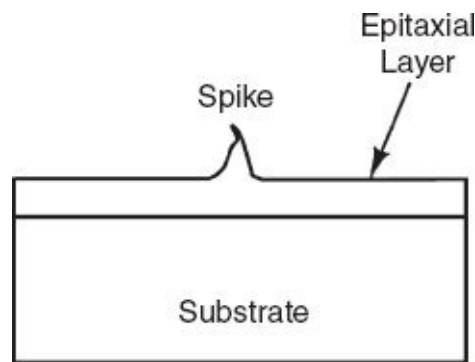


FIGURE 12.30 Epitaxial growth spike.

During the growth, a number of crystal problems can occur. One is *stacking faults*. A stacking fault is due to the inclusion of an extra atomic plane with a corresponding “dislocation” of the atoms around the plane. A stacking fault begins at the surface and “grows” to the surface of the film. The shape of the stacking fault depends on the orientation of the film and wafer. Faults in $\langle 111 \rangle$ -oriented films have a pyramidal shape (Fig. 12.31), whereas $\langle 100 \rangle$ -oriented wafers form rectangular-shaped stacking faults. The faults are detected by either X-ray or etching techniques.

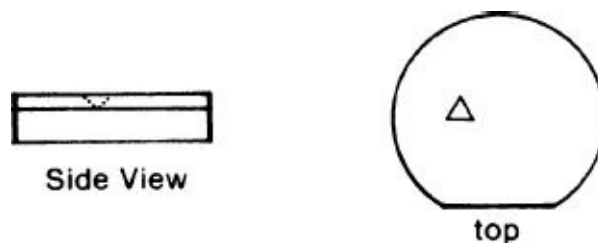


FIGURE 12.31 Stacking fault on $\langle 111 \rangle$ Si.

A growth problem associated with $\langle 111 \rangle$ wafers is *pattern shift*. This problem occurs when the deposition rate is too high and the film planes grow at an angle to the surface. Pattern shift is a problem when alignment to a subsurface pattern relies on locating it from a film surface step ([Fig. 12.32](#)). Another major growth problem is *slip*. This condition comes about from poor control of the deposition parameters and results in a “slippage” of the crystal along plane interfaces ([Fig. 12.33](#)).

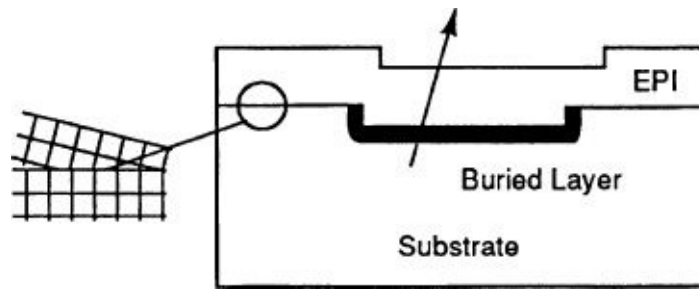


FIGURE 12.32 Epitaxial pattern shift.

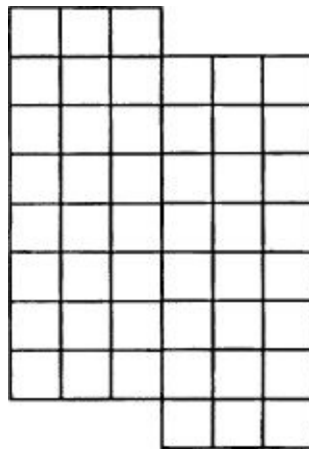


FIGURE 12.33 Crystal slip.

There are two effects that can occur during the deposition, both temperature driven: autodoping and out-diffusion. Autodoping of the growing film occurs when dopant atoms from the back of the wafer diffuse out from the wafer ([Fig. 12.34](#)), mix in the gas stream, and become incorporated into the growing film. In the film, they change the resistivity and the conductivity level. *Autodoping* in a P-type film, grown over an N-type wafer, will be less P-type than intended and have a lower P-type concentration as the autodoped atoms neutralize a number of the P-type atoms in the film.

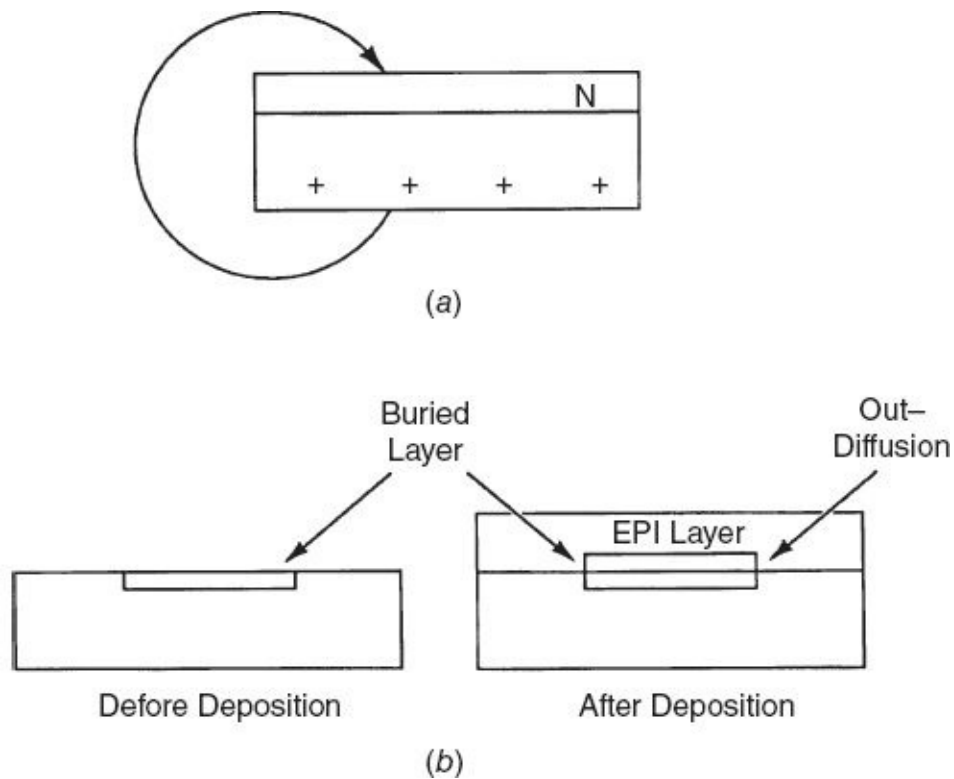


FIGURE 12.34 (a) Epitaxial autodoping and (b) out-diffusion.

Out-diffusion causes the same effect but at the epitaxial layer-wafer interface. The source of the out-diffused atoms is doped regions diffused into the wafer before the epitaxial deposition. In bipolar devices, the regions are called *buried layers* or *subcollectors*. In the usual format, the buried layer is an N-type region in a P-type wafer, over which is grown an N-type epitaxial layer. During the deposition, the N-type atoms diffuse out and become incorporated into the bottom of the epitaxial film, changing the concentration. In the extreme, the buried layer can out-diffuse up into the bipolar device structure causing electrical malfunctions.

CMOS Epitaxy

Until the late 1970s, the dominant use of epitaxial films was as the collector region of bipolar transistors. The technique provided a quality substrate for device operation and a clever means of isolating adjacent devices (see [Chap. 16](#)). A newer and perhaps more dominant use of silicon epitaxial films is for CMOS circuit wafers. The need for an epitaxial layer was driven by a CMOS circuit problem called *latch-up* (see [Chap. 16](#)). The solution is a *p-type* epi on a p^+ substrate.

Epitaxial Process

A typical epitaxial process starts with a complete and rigorous cleaning of the wafer surface prior to loading the reactor. Within the deposition chamber, a number of steps take place to correctly deposit the film. A typical SiCl_4 epitaxial process is shown in [Fig. 12.35](#). The first several steps are a gas-phase cleaning of the wafer surface. Deposition follows the cleaning with a cooldown cycle at the end. During all the steps, control of the temperature and gas flow is critical.

Cycle	Temperature	Gas	Purpose
1	Room	N_2	Purge air from system
2	Room	H_2	Reduce any organic contaminants of wafers in system
3	(Heating)	N_2	Bring system to deposition temperature
4	Deposition Temperature	HCl	Etch wafer to prepare surface for epi deposition
5	Deposition Temperature	Source + Dopant + Carrier	Grow epitaxial film
6	(Heat Off)	N_2	Purge system of reactant gases

FIGURE 12.35 Typical SiCl_4 epitaxial deposition process.

Selective Epitaxial Silicon

Advancements in epitaxial deposition systems have introduced the selective growth of epitaxial films. Whereas the epitaxial films for bipolar and CMOS substrates are deposited on the entire wafer, in selective growth they are grown through holes in either silicon dioxide or silicon nitride films. The wafer is positioned in the reactor chamber, and the epitaxial film grows directly on the silicon exposed at the bottom of the hole (Fig. 12.36). As the film grows, it takes on the crystal orientation of the underlying wafer. An advantage of such a structure is that devices formed in the surfaces of the epitaxial regions are isolated from each other by the oxide or nitride regions.

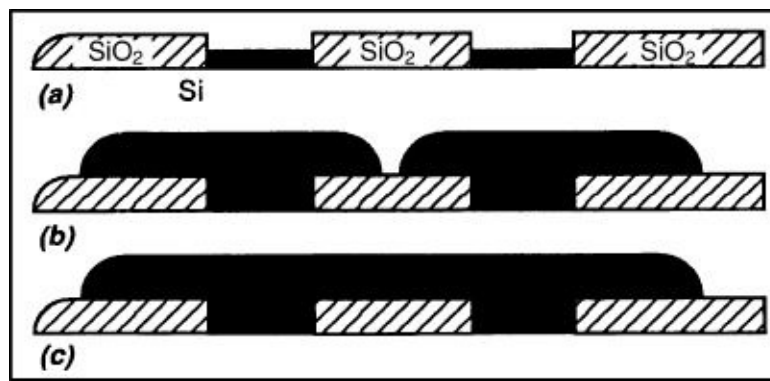


FIGURE 12.36 Steps in selective epitaxial growth.

If the deposition is allowed to continue onto the isolating surface, the structure of the film switches to a polysilicon structure. Another outcome of extended deposition is that the overlaying deposited layer becomes entirely epitaxial in nature. All of these outcomes add attractive structure options for advanced device designs.

Polysilicon and Amorphous Silicon Deposition

Until the advent of silicon-gate MOS devices ([Fig. 12.37](#)) in the mid-1970s, polysilicon layers had little or no use in device structures. Silicon-gate device technology drove the need for reliable processes to deposit thin layers of polysilicon. By the mid-1980s, polysilicon was the workhorse material of advanced devices. In addition to MOS gates, polysilicon finds use as load resistors in SRAM devices, trench fills, multilayer poly in EEPROMs, contact barrier layers, emitters in bipolar devices, and as part of silicide metallization schemes (see [Chaps. 13](#) and [16](#)).

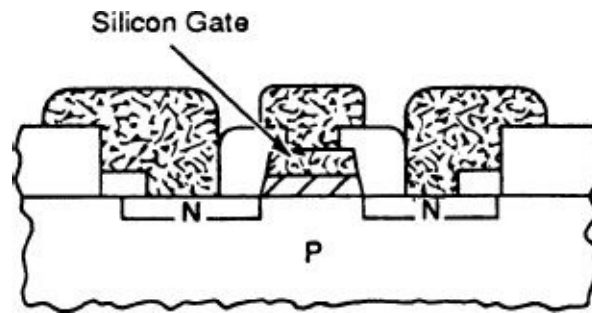


FIGURE 12.37 Cross-section of silicon gate MOS transistor.

Early processes involved simply placing the oxide-covered wafers in a horizontal APCVD system and letting the polysilicon deposit on the oxide. The major difference of early polysilicon depositions from epitaxial depositions was the use of silane sources. While silane is not favored for epitaxial film deposition, it is more than adequate for polysilicon depositions.

Typical polysilicon deposition processes take place in the 600 to 650°C range. The deposition may be from either 100 percent silane or from gas streams containing N_2 or H_2 . The structure of polysilicon was previously described as a total nonarrangement of the silicon atoms. In the case of deposited polysilicon, the structure is somewhat different. During the early stages of deposition, at temperatures below 575°C, the structure is amorphous (no structure). The polysilicon structure formed by deposition techniques consists of small pockets (crystallites or grains) of single-crystalline silicon separated by grain boundaries. This structure is called *columnar poly*.

The importance of grain size and grain boundary consistency shows up in the electrical current flow characteristics of the films. Current resistance comes as the current crosses the grain boundaries. The larger the grain boundaries, the higher the resistance. The achievement of consistent current flow from device to device and

within a device is dependent on a well-controlled polysilicon structure. One of the advantages claimed for the use of H_2 in the gas stream is the reduction of surface impurities and moisture, which in turn results in a reduced grain size. Moisture or oxygen impurities in the system cause the growth of silicon dioxide within the structure. The oxide increases the resistance of the film and its etchability in subsequent masking steps.

All of the system's usual operating parameters (temperature, silane concentration, pump speed, nitrogen flow, and other gas flows²⁴) affect the deposition rate and the grain size. Often, the wafers will receive a postdeposition anneal in the 600°C range to further crystallize the film. The process of recrystallization goes on whenever the wafers go through a high-temperature process. The grain size and electrical parameters of the polysilicon film on the finished device or circuit are never the same as the deposited film.

Also influencing the grain size is the presence of dopants in the gas stream. In many devices or circuits, a strip of polysilicon functions as a conductor which requires doping to decrease its resistivity. Doping can be done by diffusion before or implantation after the deposition.

In situ doping takes place by adding gas dopant sources in the source cabinet and metering them into the chamber. When diborane (boron source) is added, there is a large increase in the deposition rate. An opposite effect takes place when phosphine (phosphorus source) or arsine (arsenic source) is the dopant gas. Undesirable effects of *in situ* doping are a loss of film uniformity, doping uniformity, and control of the deposition rate. Doped polysilicon film resistivities are lower than those of equally doped epitaxial or bulk silicon. The lower resistivities are due to dopants being trapped in the grain boundaries.

Most polysilicon layers are deposited with LPCVD systems that provided good productivity and lower deposition temperatures. LPCVD provides good step coverage ([Fig. 12.38](#)), a requirement because polysilicon layers are usually deposited later in the process, and the surface has become varied in its topography. Single-chamber polysilicon LPCVD systems offer the advantage of higher deposition rates without raising the temperature.²⁵

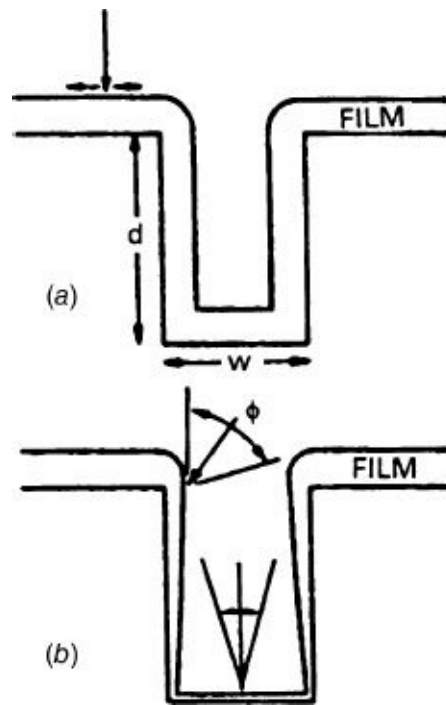


FIGURE 12.38 Step coverage: (a) good step coverage and (b) nonconformal coverage.

SOS and SOI

These two acronyms stand for *silicon on sapphire* and *silicon on insulator*. Both refer to the deposition of silicon on a non-semiconductor surface. The need for such structures came about from the limits placed on some MOS devices by the presence of a semiconducting substrate under the active device. These problems are resolved by forming a silicon layer on an insulating substrate. The first combination for this purpose was silicon on sapphire. As different substrates were investigated, the term was expanded to the more general, silicon on insulator.

One technique is a direct deposition on the substrate followed by a recrystallization process (laser heating, strip heaters, oxygen implantation) to create a usable film.²⁶ Another approach is a selective deposition through holes in a surface oxide with an overgrowth to form the continuous film.

Another SOI method is *SIMOX*. In this process, the top layer of a wafer is converted to oxide with a heavy oxygen implant. An epitaxial layer is grown on top of the oxide. There is some exploration of *bonded wafers*. This approach has two wafers bonded together followed by the thinning (grinding and polishing) of one to device layer thickness.²⁷

Gallium Arsenide on Silicon

Gallium arsenide is a great III–V semiconductor material. However, it is fragile and limited to wafer sizes of 4-in (102-mm) diameter. Attempts to grow GaAs films on silicon wafers have been thwarted by the mismatch of the lattice sizes of each. During the growing process, the mismatch causes dislocations that degrade device performance. A new approach is to first grow a thin layer of strontium titanate on the silicon wafer. It reacts with the silicon to form an amorphous silicon dioxide layer. When a film of GaAs is deposited, the silicon dioxide layer acts as a cushion that absorbs the mismatch, thus allowing a single crystalline layer.²⁸

Insulators and Dielectrics

CVD is the favored method of depositing films that will function in the device or circuit as insulators or dielectrics. The two films in widespread use are silicon dioxide and silicon nitride. In general, the two films find a multiplicity of uses in device and circuit designs. While they have processing and quality differences, they must meet the same general requirements as other deposited films.

Silicon Dioxide

Deposited silicon dioxide films are best known from their long-term use as a final passivation layer covering the completed wafer. In this role, they provide physical and chemical protection to the underlying circuit devices and components. Deposited silicon dioxide films used as a protective top layer are known by the proprietary terms *Vapox*, *Pyrox*, and *Silox*[®]. *Vapox* (vapor-deposited oxide) is a term coined by Fairchild engineers. *Pyrox* stands for pyrolitic oxide. *Silox* is a registered trademark of Applied Materials, Inc. Sometimes, the layer is simply called a *glass*. This protective role has expanded, and deposited silicon oxide layers are used as interdielectric layers in multimetallization schemes, as insulation between polysilicon and metallization layers, as doping barriers, as diffusion sources, and as isolation regions. Silicon dioxide has become a major part of silicon gate structures.

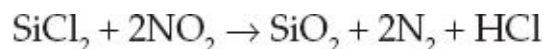
There are gate stacks, consisting of thermal oxide or silicon dioxide or oxynitride/silicon dioxide (TEOS deposited), and various silicon dioxide fillers for plugs in multimetall designs.²⁹

CVD-deposited silicon dioxide films vary in structure and stoichiometry from thermally grown oxides. Depending on the deposition temperature, deposited oxides will have a lower density and different mechanical properties, such as index of refraction, resistance to cracking, dielectric strength, and etch rate. These factors are highly affected by the addition of dopants to the film. In many processes, the deposited film will receive a high-temperature anneal, a process called *densification*. After the densification, the deposited silicon dioxide film is close to the structure and properties of a thermal oxide.

The need for a low-temperature-deposited SiO₂ was dictated by the unacceptable alloying of aluminum and silicon at temperatures above 450°C. The early deposition process used was a horizontal conduction heated APCVD system from silane and oxygen by the reaction $\text{SiH}_4 + \text{O}_2 \rightarrow \text{SiO}_2 + 2\text{H}_2$

This process produced films that were of unacceptable quality for use in advanced-device designs and on larger wafers due to the poorer film quality produced by the 450°C deposition temperature.

The development of LPCVD systems made possible higher-quality films, especially for the factors of step coverage and lower stress. LPCVD processes are the preferred deposition techniques from both quality and productivity considerations. High-temperature (900°C) LPCVD of silicon dioxide is performed with a dichlorosilane reaction with nitrous oxide.



Tetraethyl Orthosilicate

By far, the majority of silicon dioxide are deposited from $\text{Si}(\text{OC}_2\text{H}_5)_4$ sources. The source is known as tetraethyl orthosilicate (TEOS). TEOS history goes back to the 1960s. Early systems relied on the simple pyrolysis of the TEOS in the 750°C range. Current depositions are based on the hot-wall LPCVD systems established in the 1970s, with temperatures in the 400°C+ range. TEOS sources used with plasma-assistance (PECVD or PETEOS) allowed deposition temperatures in the sub-400°C range.³⁰ This process faces limits on conformal coverage of high-aspect ratio patterns in 0.5-μm devices. Step coverage is improved by the addition of ozone (O_3) to the gas stream.³¹

Another option is the reaction of silane with nitrous oxide in argon plasma.



Doped Silicon Dioxide

Silicon dioxide layers are doped to improve their protective characteristics and flow properties, or for use as dopant sources. The earliest dopant used with deposited oxides was phosphorus. The phosphorus source is phosphine (PH_3) gas added to the deposition gas stream. The resultant glass is called *phosphorus silicate glass* (PSG). Within the glass, the phosphorus is in the form of phosphorus pentoxide (P_2O_5), making the glass a dual compound or, more correctly, a binary glass.

The role of the phosphorus is threefold. The added dopant increases the moisture-barrier property of the glass. Mobile ionic contaminants become attached to the phosphorus and are prevented from traveling into the wafer surface. This action is called *gettering*. The third result is an increase of the flow characteristics ([Fig. 12.39](#)), which aid the planarization of the glass surface after a heating step in the 1000°C range. The phosphorus content is limited to about 8 weight by percent. Above this level, the glass becomes hygroscopic and attracts moisture. The moisture can react with the phosphorus, form phosphoric acid, and attack underlying metal lines.

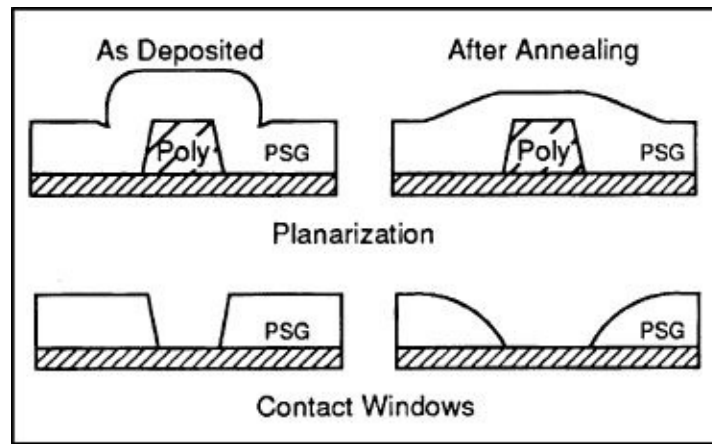


FIGURE 12.39 Planarization of surface by flowing glass.

Boron is often added to the glass from a diborane (B_2H_6) source. The purpose of the boron is to also aid the flow characteristics ([Fig. 12.39](#)). The resultant glass is called a borosilicate glass (BSG). The boron and phosphorus are often used together in the glass. The result is referred to as borophosphorus silicate glass (BPSG).

Silicon Nitride

Silicon nitride is a replacement for silicon dioxide uses, especially for top-layer protection. Silicon nitride is harder, which provides better scratch protection, is a better moisture and sodium barrier (without doping), has a higher dielectric strength, and resists oxidation. The latter property has led to its use in the local oxidation of silicon (LOCOS) for isolation purposes. [Figure 12.40](#) illustrates the process, where patterned islands of silicon nitride prevent oxidation under the islands. After thermal oxidation and removal of the nitride, there are wafer surface regions ready for device formation, separated by isolating regions of oxide. A disadvantage of silicon nitride is that it does not flow as easily as silicon oxide and is more difficult to etch. The etch restriction has been overcome with the development of plasma-etch processes.

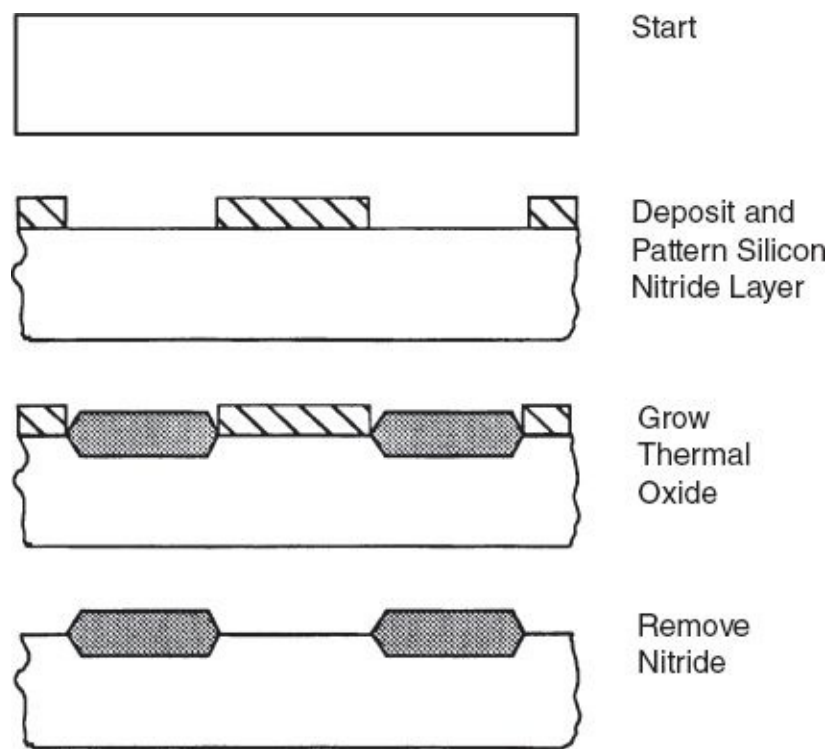


FIGURE 12.40 LOCOS process.

An early limit on the use of silicon nitride protective films was the lack of a low-temperature deposition process. In APCVD systems, a temperature of 700 to 900°C is required for the deposition of silicon nitride from silane or dichlorosilane ([Fig. 12.41](#)). The result is a film with the composition Si_3N_4 . The reactions also take place in LPCVD reactors but at a temperature low enough for deposition over an aluminum metallization layer. The advent of PECVD has opened up the use of

different source chemistries. One use is silane reacted with ammonia (NH₃) or nitrogen in the presence of an argon plasma.

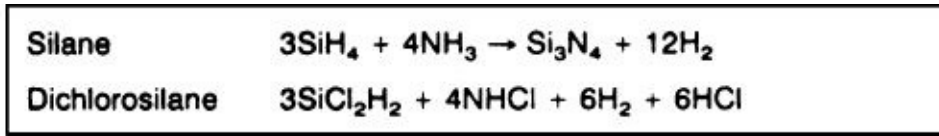


FIGURE 12.41 Silicon nitride deposition reactions.

High-k and Low-k Dielectrics

In addition to the dielectric films mentioned above, there are numerous other types deposited for special applications. They fall into the broad categories of high-k and low-k. In [Chap. 2](#), the basics of a capacitor were described. For review, the k value of a material is called its *dielectric constant*. It relates to the level of capacitance a capacitor will have. High-k materials produce capacitors with high capacitance that are desirable for charge storage. They also have been incorporated into MOS devices as the gate dielectric (gate stacks). ALD and MOCVD deposition systems are favored for this use due to their high level of thickness control.³² Uses of high-k dielectrics are discussed in [Chap. 16](#).

Low-k materials are used in metallization systems as dielectric barriers in metallization systems on the wafer. The introduction of copper interconnects has particularly driven the use of low-k materials. They are applied by PECVD techniques and by spinon applications. These low-k materials are discussed in [Chap. 13](#).

Conductors

The traditional metal conductors of aluminum and aluminum alloys are deposited by evaporation or sputtering techniques. The advent of silicon gate MOS transistors added doped polysilicon as a device conductor. Plus, the advent of multimetal structures and new conducting materials has thrust CVD and PVD techniques into the conducting metal business. The techniques and use of these deposited metals are explained in the next chapter.

Review Topics

Upon completion of this chapter, you should be able to:

1. Name the parts of a CVD reactor.
 2. Describe the principle of chemical vapor deposition.
 3. List the conductor, semiconductor, and insulator materials deposited by CVD techniques.
 4. Know the difference between atmospheric CVD, LPCVD, hot-wall, and cold-wall systems.
 5. Explain the difference between epitaxial and polysilicon layers.
-

References

1. Hayes, J. and Van Zant, P., *CVD Today Seminar Manual*, Semiconductor Services, 1985, San Jose, CA:9.
2. Wolf, S. and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Sunset Beach, CA:165.
3. Hammond, M. L., "Epitaxial Silicon Reactor Technology: A Review," *Solid State Technology*, May 1988:160.
4. Ibid.
5. Hayes, J. and Van Zant, P., *CVD Today Seminar Manual*, 1985, Semiconductor Services, San Jose, CA:13.
6. Wolf, S. and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Sunset Beach, CA:165.
7. Meyerson, B., Kaiser, H., and Schultz, S., "Extending Silicon's Horizon through UHV/CVD," *Semiconductor International*, Cahners Publishing, Mar. 1994:73.
8. Singer, P., "The Future of Dielectric CVD: High Density Plasma?" *Semiconductor International*, Jul. 1997:127.
9. Applied Materials, Centura Isprint Tungston ALD/CVD Product Description, <http://www.appliedmaterials.com/technologies/library/centura-isprint-tungsten-aldcvd>. June 2013.
10. Braum, A., "ALD Breaks Materials, Conformity Barriers," *Semiconductor International*, Oct. 2001:52.
11. Williams, R., *Gallium Arsenide Processing Techniques*, 1984, Artech House, Dedham, MA:44.
12. Wolf, S. and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Sunset Beach, CA:157.
13. Singer, P., "Molecular Beam Epitaxy," *Semiconductor International*, Oct. 1986:42.
14. Ibid.
15. Ibid.
16. Burggraaf, P., "The Growing Importance of MOCVD," *Semiconductor International*, Nov. 1986:47.
17. Department of Electronic Materials Engineering, *Tech Note*, 2002, The Australian National University.
18. Burggraaf, P., "The Growing Importance of MOCVD," *Semiconductor International*, Nov. 1986:48.

- [19.](#) Burggraaf, P., "The Status of MOCVD Technology," *Semiconductor International*, Cahners Publishing, Jul. 1993:81.
- [20.](#) Thompson, A., Stall, R., and Droll, B., "Advances in Epitaxial Deposition Technology," *Semiconductor International*, Cahners Publishing, Jul. 1994:173.
- [21.](#) Department of Electronic Materials Engineering, *Tech Note*, 2002, The Australian National University.
- [22.](#) Hammond, M. L., "Epitaxial Silicon Reactor Technology: A Review," *Solid State Technology*, May 1988:159.
- [23.](#) Roberge, R. P., "Gaseous Impurity Effects in Silicon Epitaxy," *Semiconductor International*, Jan. 1988:81.
- [24.](#) Sze, S. M., *VLSI Technology*, 1983, McGraw-Hill, New York, NY:119.
- [25.](#) Venkatesan, M. and Beinglass, I., "Single-Wafer Deposition of Polycrystalline Silicon," *Solid State Technology*, PennWell Publishing, Mar. 1993:49.
- [26.](#) Jastrzebski, L., "Silicon CVD for SOI: Principles and Possible Applications," *Solid State Technology*, Sep. 1984:239.
- [27.](#) Yallup, K., "SOI Provides Total Dielectric Isolation," *Semiconductor International*, Cahners Publishing, Jul. 1993:134.
- [28.](#) Singer, P., "GA-As-on-Silicon, Finally!" *Semiconductor International*, October, 2001:36
- [29.](#) Singer, P., "Directions in Dielectrics in CMOS and DRAMs," *Semiconductor International*, Cahners Publishing, Apr. 1994:57.
- [30.](#) Chin, B. L. and van de Ven, E. P., "Plasma TEOS Process for Interlayer Dielectric Applications," *Solid State Technology*, Apr. 1988:119.
- [31.](#) Maeda, K. and Fisher, S., "CVD TEOS/O₃: Development History and Applications," *Solid State Technology*, PennWell Publishing, Jun. 1993:83.
- [32.](#) Deweerdt, W., Delabie, A., Van Elshocht, S., et al., "ALD vs. MOCVD for High-k Deposition in 45-nm CMOS and Below," *Future Fab Intl.*, Jan. 2006:20.

CHAPTER 13

Metallization

Introduction

Fabrication of circuits is divided into two major segments. First the active and passive parts of the components are fabricated in and on the wafer surface. This is called the *front end of the line* (FEOL). Second is the metal system necessary to connect the devices and different layers added to the chip. This is called the *back end of the line* (BEOL). In this chapter, the materials, specifications, and methods used to complete the metallization segment are presented along with other uses of metals in chip manufacturing. Vacuum pumps [used in chemical vapor deposition (CVD), ion implant, evaporation, and sputtering systems] are explained at the end of the chapter.

The most common use of metal films in semiconductor technology is for surface wiring. The materials, methods, and processes of “wiring” the component parts together is generally referred to as the *metallization process*. Depending on device complexity and performance requirements, the circuit may require a single-metal system, or a multiple-level system. It may use an aluminum alloy or copper as the conducting metal.

Deposition Methods

• Sputtering of aluminum alloys and other metals • Low-pressure CVD of polysilicon, tungsten, and other refractory metals • Dual-damascene copper processes with electroplating

Metallization techniques, like other fabrication processes, have undergone improvements and evolution in response to new circuit requirements and new materials. The mainstay of metal deposition up to the mid-1970s was vacuum evaporation of aluminum, gold, and fuse metals for programmable read only memory (PROM) devices. The advent of multilayer metal systems and alloys, along with the need for better step coverage, led to the introduction of sputtering as the standard deposition technique for very large-scale integration (VLSI) circuit fabrication. Refractory metal use has added the second technique, CVD, to the arsenal of the metallization engineer. Copper was introduced as a primary metal with the development of the dual-damascene process with electroplating.

Multilayer systems led to the development of barrier and adhesion layers, plugs, and intermediate dielectric layers. The basics of single-layer metal and multilayer metal systems are explored below.

Single-Layer Metal Systems

In the MSI era metallization was relatively straightforward ([Fig. 13.1](#)), requiring

only a single-level metal process. Small holes, called *contact holes* or *contacts*, are etched through the surface layers, down to the contact areas on the individual devices. These are created in a photomasking step called contact masking. Following contact masking, a thin layer (currently about $0.5\ \mu\text{m}$) of the conducting metal is deposited by vacuum evaporation, sputtering, or CVD techniques over the entire wafer. The unwanted portions of this layer are removed by a conventional photomasking and etch procedure or by liftoff. This step leaves the surface covered with thin lines of the metal that are called *leads*, *metal lines*, or *interconnects*. Generally, a heat-treatment step, called *alloying*, is performed after metal patterning to ensure good electrical contact between the metal and the wafer surface. This basic process is shown in [Fig. 13.1](#).

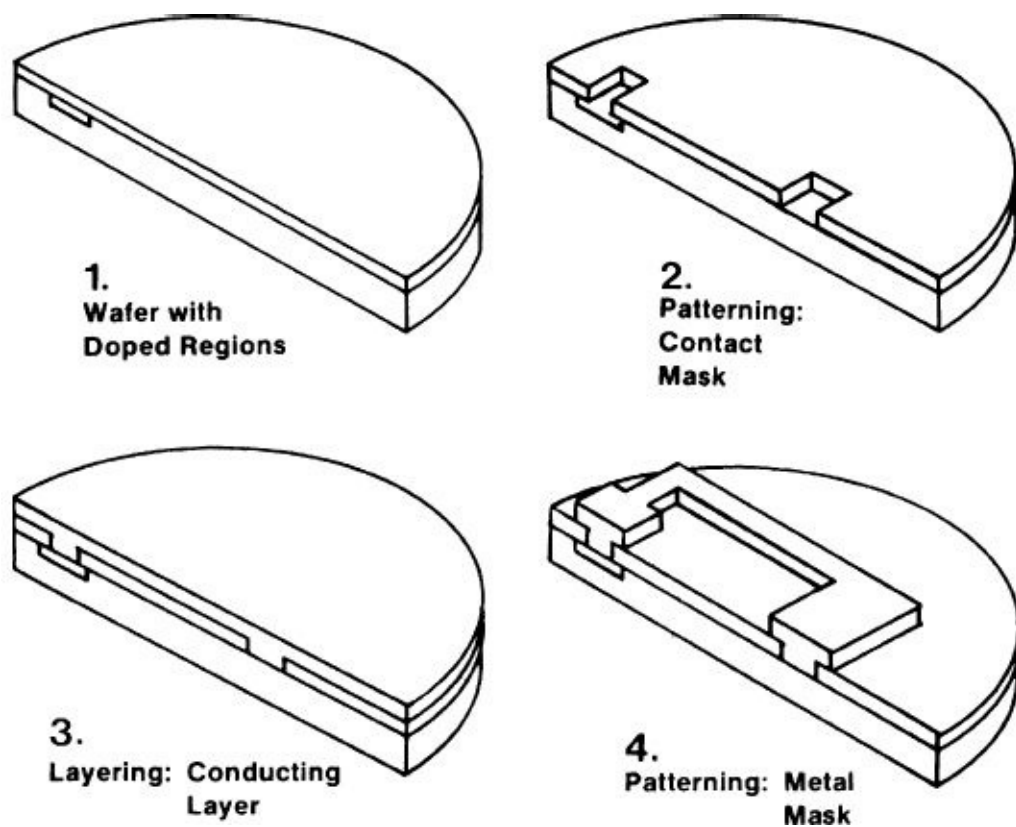


FIGURE 13.1 First-level metallization sequence.

Regardless of the structure, a metal system must meet the following criteria: • Good electrical current-carrying capability (current density) • Good adhesion to the top surface of the wafer (usually SiO_2) • Ease of patterning

- Good electrical contact with the wafer material • High purity
- Corrosion resistance • Long-term stability
- Capable of deposition in uniform layers that are void and hillock-free films • Uniform grain structure

Multilevel Metal Schemes

Increasing chip density has placed more components on the wafer surface, which in turn has *decreased* the area available for surface wiring. The answer to this dilemma has been multilevel metallization schemes ([Fig. 13.2](#)). By 2120, the ITRS is projecting 15 to 20 metal layers.¹ A basic two-metal stack is shown in [Fig. 13.3](#). The stack starts with a *barrier layer* formed by silicidation of the silicon surface to produce a lowered electrical resistance between the surface and the next layer. Barrier layers also prevent alloying of aluminum and silicon if pure aluminum is the conducting material. Next comes a layer of dielectric material, called an *intermetallic dielectric layer* (IDL or IMD) that provides electrical isolation between metal layers. This dielectric may be a deposited oxide, silicon nitride, or a polyimide film. This layer receives a masking step that etches new contact holes, called *vias* or *plugs*, down to the first-level metal. Conducting plugs are created by depositing conducting material into the hole. Next, the second-level metal layer is deposited and patterned. The IMD/plug/metal deposition or patterning sequence is repeated if there are subsequent layers. A multilevel metal system is more costly, of lower yield, and requires greater attention to planarization of the wafer surface and intermediate layers to create good current-carrying leads.

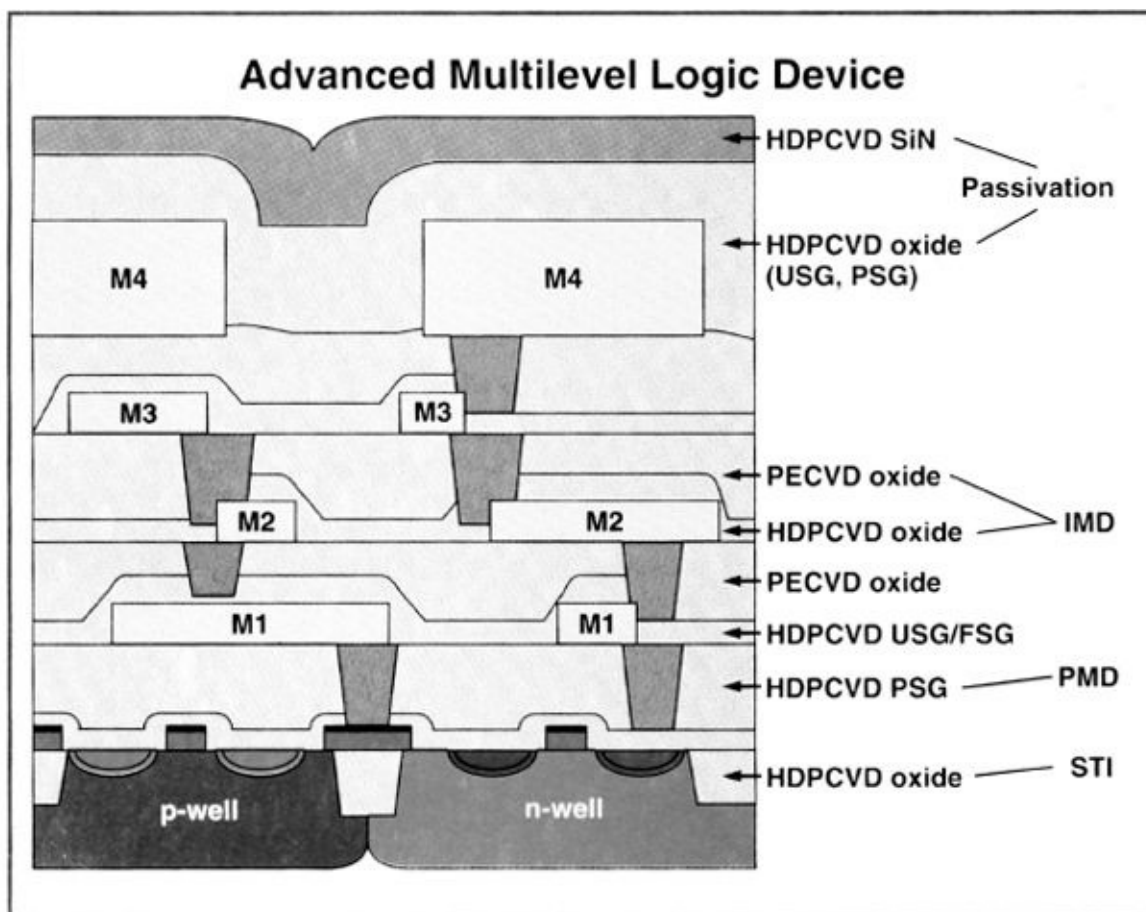


FIGURE 13.2 Multimetall-level structure. (Courtesy of Semiconductor International, July 1997.)

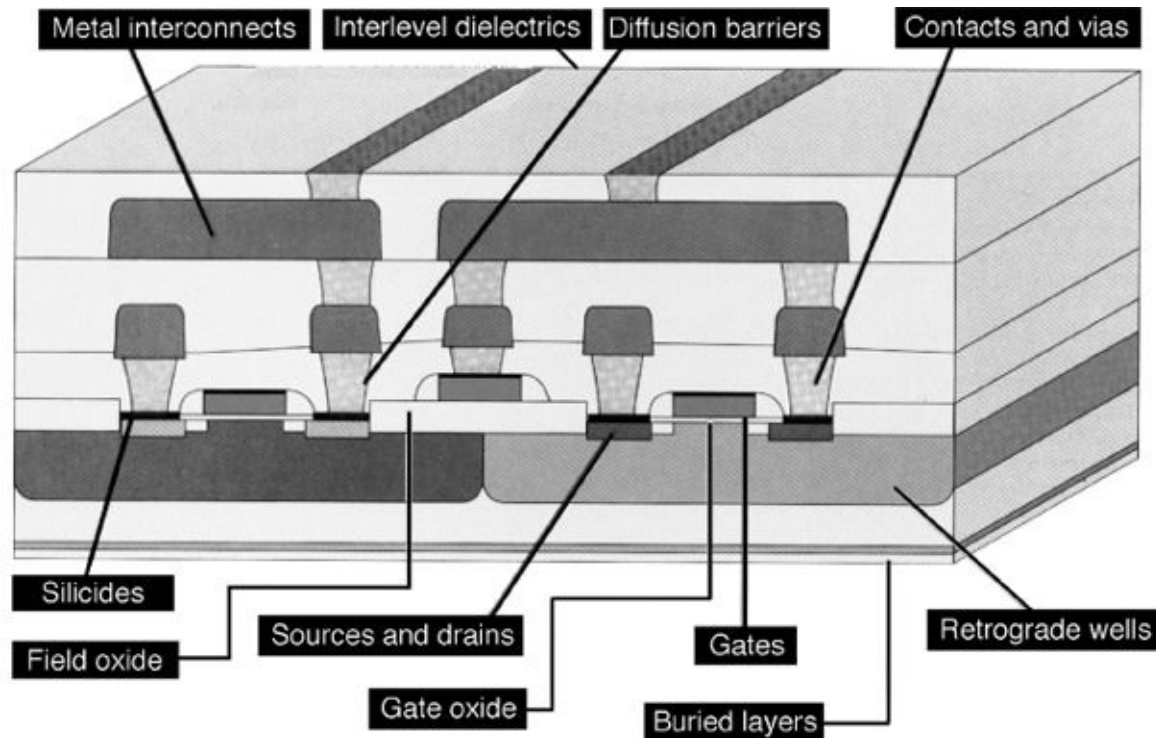


FIGURE 13.3 Two-metal structure. (Courtesy of Semiconductor International, January 1998.)

Conductor Materials

Aluminum

This section addresses the three primary materials used for the metal interconnection layers. Prior to the development of VLSI-level circuits, the primary metallization material was pure aluminum. The choice of aluminum and its limitations are instructive to the understanding of metallization systems in general. From an electrical conduction standpoint, aluminum is less conductive than copper and gold. Copper, if used as a direct replacement for aluminum, has a high contact resistance with silicon and raises havoc with device performance if it gets into the device areas. Aluminum emerged as the preferred metal because it avoids the problems just mentioned. It has a low enough resistivity ($2.7 \mu\Omega\text{-cm}$),² and good current-carrying density. It has superior adhesion to silicon dioxide, is available in high purity, has a naturally low contact resistance with silicon, and is relatively easy to pattern with conventional photolithography processes. Aluminum sources are purified to 5 to 6 “nines” of purity (99.999 to 99.9999 percent).

Aluminum-Silicon Alloys

Shallow junctions in the wafer surface presented one of the first problems with the use of pure aluminum leads. The problem came with the need to bake aluminum-silicon interfaces to stabilize the electrical contact. This type of contact is called *ohmic* because the voltage-current characteristics behave according to Ohm's law. Unfortunately, aluminum and silicon dissolve into each other and, at 577°C, reach a eutectic formation point. A eutectic formation occurs when two materials heated in contact with each other melt at temperatures much lower than their individual melting temperatures. Eutectic formations occur over a temperature range, and the aluminum-silicon eutectic starts to form at about 450°C, also the temperature necessary for good electrical contact. The problem (often called *spiking*) is acute with shallow junctions. If the alloy region is deep, it can extend completely through the junction, shorting it out ([Fig. 13.4](#)).

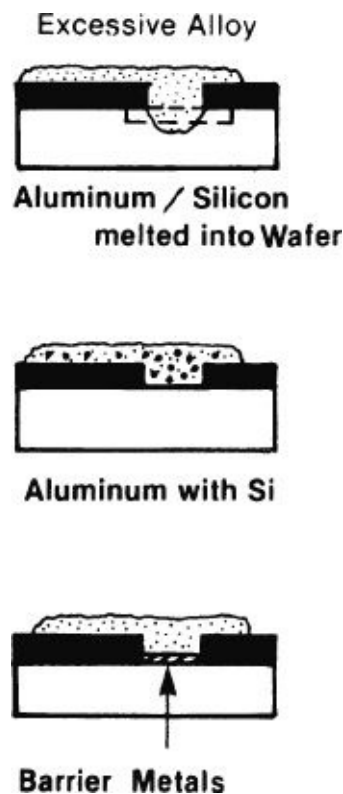


FIGURE 13.4 Eutectic alloying of aluminum and silicon contacts.

There are two solutions to this. One is a barrier metal layer (see the section on barrier metals) that separates the aluminum and silicon and prevents the eutectic alloy from forming. The second is an alloy of aluminum with 1 to 2 percent silicon. During the contact-heating step, the aluminum alloys more with the silicon in the alloy and less with the silicon from the wafer. This process is not 100 percent effective, and some alloying between the aluminum and wafer always occurs.

Aluminum-Copper Alloys

Aluminum suffers a mechanism called *electromigration*. The problem occurs when long, skinny leads of aluminum carry high currents over long distances, as is the situation in VLSI/ultra-large-scale integrated (ULSI) circuits. The current sets up an electric field in the lead that is higher at the input side of the lead and decreases along the lead to the output contact. Also, heat generated by the flowing current sets up a thermal gradient along the lead. The aluminum along the lead becomes mobile and diffuses within itself along the direction of the two gradients. The first effect is thinning of the lead. In the extreme, the lead can become completely separated. Unfortunately, this event usually takes place after the circuit is in operation in the field, causing a failure of the chip. Prevention or moderation of electromigration is achieved by depositing an alloy of aluminum and 0.5 to 4 percent³ copper or an alloy of aluminum and 0.1 to 0.5 percent titanium. Aluminum alloys containing both copper and silicon are often sputter deposited on the wafer to resolve both alloying and electromigration problems.

The early deposition of aluminum alloys was by putting separated sources in an evaporation system. This leads to increased complexity for the deposition equipment and process. Also the aluminum alloy films have a higher resistivity compared to pure aluminum. The amount of the increase varies with the alloy composition and heat treatments but can be as much as 25 to 30 percent.⁴

Barrier Metals

A method of preventing the eutectic alloying of silicon and aluminum metallization is by using a barrier layer. Both titanium-tungsten (TiW) and titanium nitride (TiN) layers are used. TiW is sputter-deposited onto the wafer into the open contacts before the aluminum or aluminum alloy deposition takes place. The TiW deposited on the field oxide is removed from the surface during the aluminum etch step. Sometimes, a first layer of platinum silicide is formed on the exposed silicon before the TiW is deposited.

Titanium nitride layers can be placed on the wafer by all the deposition techniques: evaporation, sputtering, and CVD. It can also be formed by the thermal nitridation of a titanium layer at 600°C in an N₂ or NH₃ atmosphere.⁵ CVD titanium nitride layers have good step coverage and can fill submicron contacts. A layer of titanium is required under TiN films to provide a high-conductivity intermediate with silicon substrates.

With copper metallization, the barrier is also critical. Copper inside the silicon ruins device performance. Barrier metals used with copper metalization are titanium nitride (TiN), tantalum (Ta), and tantalum nitride (TaN).⁶

Refractory Metals and Refractory Metal Silicides

Although the limitations of electromigration and eutectic alloying have been made manageable by aluminum alloys and barrier metals, the issue of contact resistance may prove to be the final limit on aluminum metallization. The overall effectiveness of a metal system is governed by the resistivity, length, thickness, and total contact resistance of *all* the metal-wafer interconnects. In a simple aluminum system, there are two contacts: the silicon-aluminum interconnect and the aluminum interconnect-bonding wire. In a ULSI circuit with multilevel metal layers, barrier layers, plug fills, polysilicon gates and conductors, and other intermediate conductive layers, the number of connections becomes very large. The addition of all the individual contact resistances can dominate the conductivity of the metal system ([Fig. 13.2](#)).

Contact resistance is influenced by the materials, the substrate doping, and the contact dimensions. The smaller the contact size, the higher the resistance. Unfortunately, ULSI chips require smaller contact openings, and a large gate array chip surface can be as much as 80 percent contact area.² These two factors make the contact resistance a dominant factor in VLSI metal system performance. Aluminum-silicon contact resistance, along with the alloying problem, has led to the investigation of other metals for VLSI or ULSI metallization. Polysilicon has a lower contact resistance than aluminum and is in use in metal oxide semiconductor (MOS) circuits ([Fig. 13.5](#)). This is the legendary *silicon gate* MOS device structure.

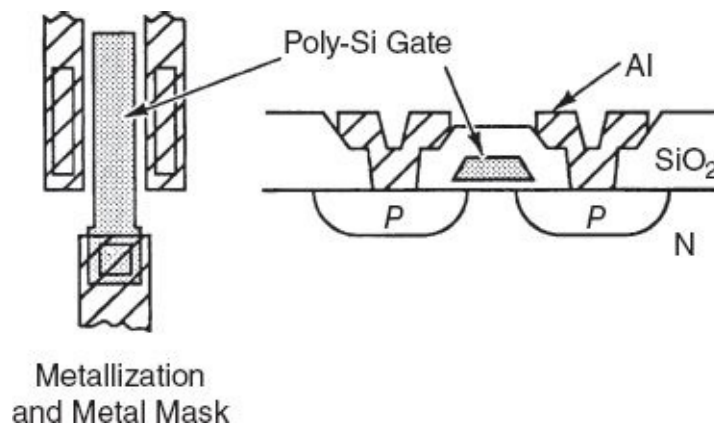


FIGURE 13.5 Silicon-gate electrode extended for metallization lead.

Refractory metals and their silicides offer lower contact resistance. The refractory metals of interest are titanium (Ti), tungsten (W), tantalum (Ta), and molybdenum (Mo). Their silicides form when they are alloyed on a silicon surface (WSi_2 , TaSi_2 , MoSi_2 , and TiSi_2). The refractory metals were first proposed for metallization in the 1950s, but they stayed in the background due to a lack of a reliable deposition method. That situation has changed with the development of low pressure CVD (LPCVD) and sputtering processes.

All modern circuit designs, especially MOS circuits, use refractory metals or their silicides as intermediate (plugs), barriers, or conducting layers. The lower resistivities and lower contact resistances ([Fig. 13.6](#)) make them attractive for conducting films, but impurities and deposition uniformity problems make them less attractive for MOS gate electrodes. The solution to the problem has been the polycide and silicide gate structures, which are combinations of a silicon gate topped by a silicide. The details of this structure are explained in [Chap. 16](#).

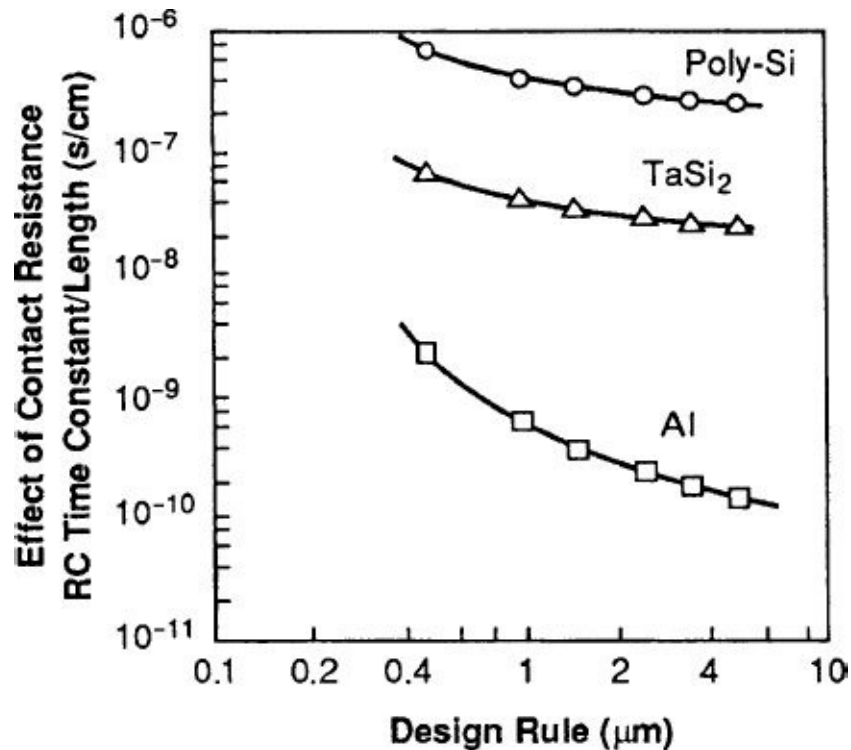


FIGURE 13.6 Effect of contact resistance on resistance capacitance (RC) time constant.

Plugs

A popular use of refractory metals is the filling of via holes in multilevel metal structures. The process is called *plug filling*, and the filled via is called a *plug* ([Fig. 13.7](#)). While polysilicon and aluminum have been used for plugs, tungsten (W from the other name: wolfram). The vias are filled by either selective tungsten deposition through surface holes onto the first layer metal or by, the standard, blanket CVD techniques.⁸ Of the available refractory metals, tungsten finds a lot of use as aluminum-silicon barriers, MOS gate interconnects, and for via plugs. Tungsten is favored for its superior step coverage, lowered electrical resistance, resistance to electromigration, and high-temperature tolerance. However, its contact resistance with silicon and adhesion challenges requires additional layers that form the classic tungsten *stack*. Tin is deposited first (contact) followed by TiN (adhesion) before the tungsten deposition. Additionally, the via can be overfilled with tungsten and back-

etched or flattened by a chemical-mechanical processing (CMP) process.

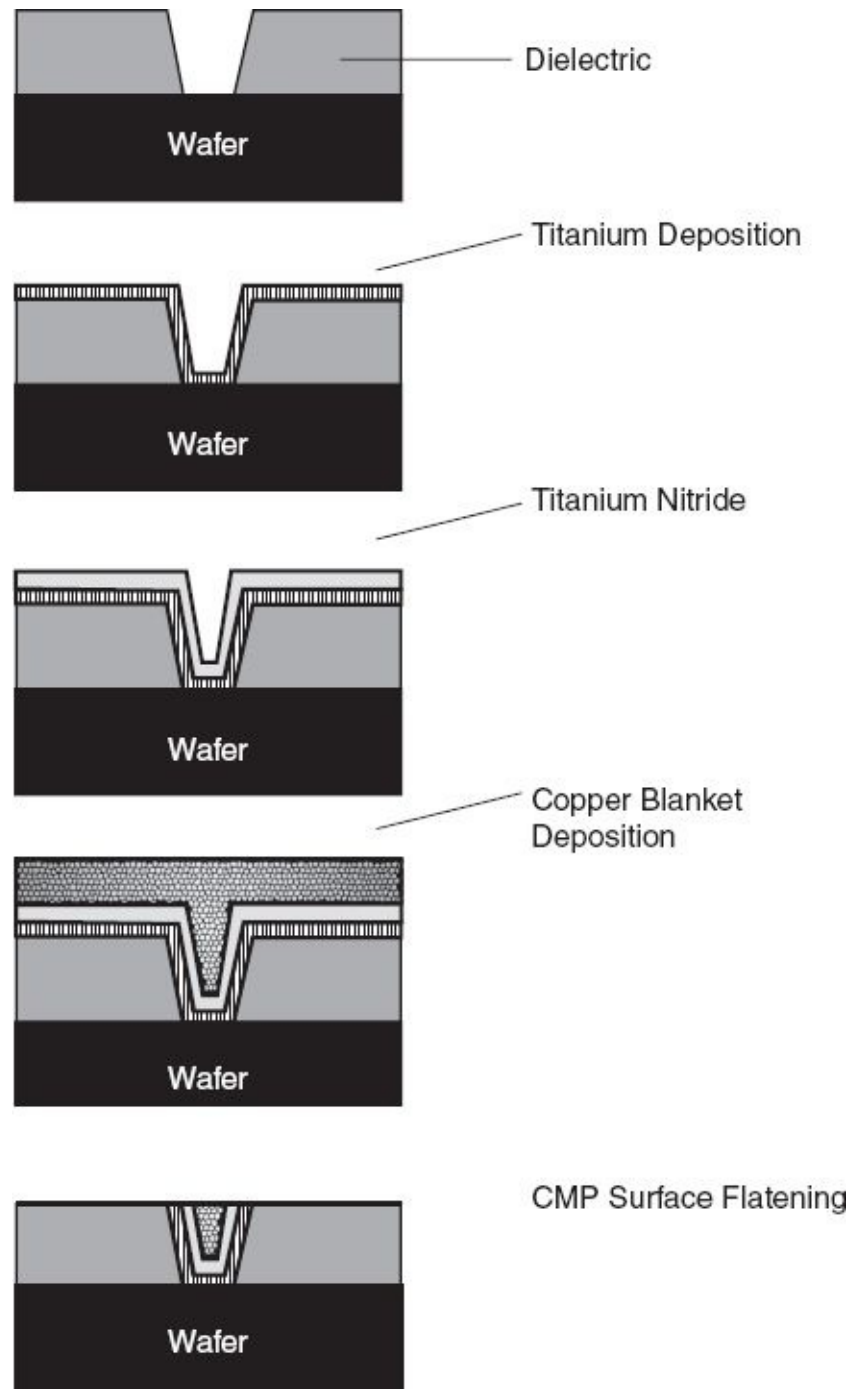


FIGURE 13.7 Tungsten plug process steps.

Sputter Deposition

The historic metal deposition process was vacuum evaporation.⁹ It took place in a stainless steel bell jar with wafers held in rotating domes over a metal source heated to evaporation levels by an electron stream ([Fig. 13.8](#)). Its limitations were met by

the introduction of aluminum alloys and step coverage into high aspect via holes. Different metals evaporate at different rates that made depositing uniform alloys difficult. And the advent of larger wafer diameters limited production rates in evaporation systems. Sputter deposition (sputtering) solved these problems and is the standard metal-deposition method.

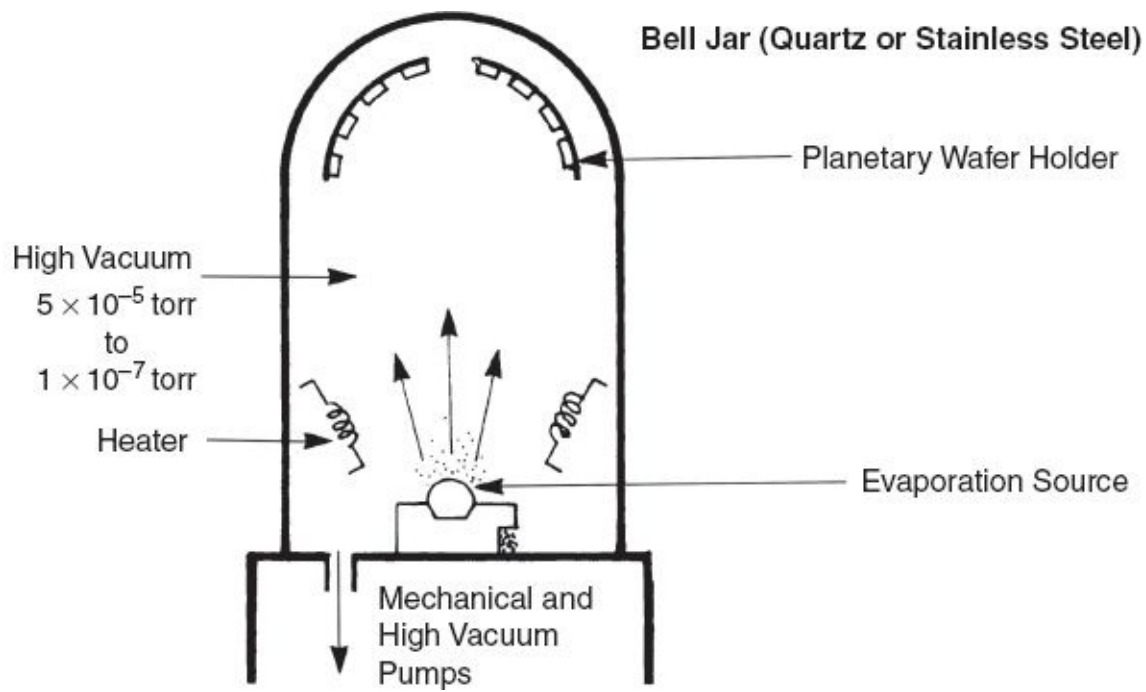


FIGURE 13.8 Vacuum evaporator.

Sputter deposition is a physical vapor deposition (PVD) process, first formulated in 1852 by Sir William Robert Grove.¹⁰ Sputtering (in general) can deposit any material on any substrate. It is widely used to coat costume jewelry and put optical coatings on lenses and glasses. Discussion of the benefits of sputtering to the semiconductor industry is best left until the principles and methods of sputtering have been covered.

Inside the vacuum chamber is a solid slab, called a *target*, of the desired film material to be deposited ([Fig. 13.9](#)). The target is electrically grounded. Argon gas is introduced into the chamber and is ionized to a positive charge. The positively charged argon atoms are attracted to the grounded target and accelerate toward it. During the acceleration, they gain momentum, which is force, and strike the target. At the target, a phenomenon called *momentum transfer* takes place. Just as a cue ball transfers its energy to the other balls on a pool table, causing them to scatter, the argon ions strike the slab of film material, causing its atoms to scatter ([Fig. 13.10](#)). This is a physical process. Literally, the argon atoms “knock off” atoms and molecules from the target, sending them into the chamber. This is the sputtering

activity. The sputtered atoms or molecules scatter in the chamber with some coming to rest on the wafer. A principal feature of a sputtering process is that the target material is deposited on the wafer without chemical or compositional change.

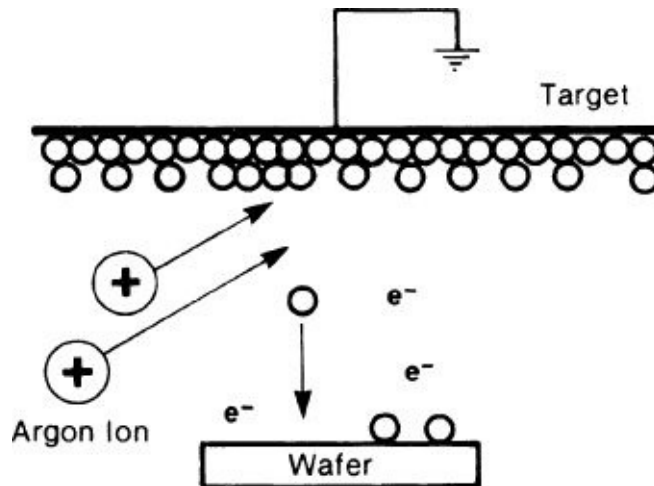


FIGURE 13.9 Principle of sputtering.

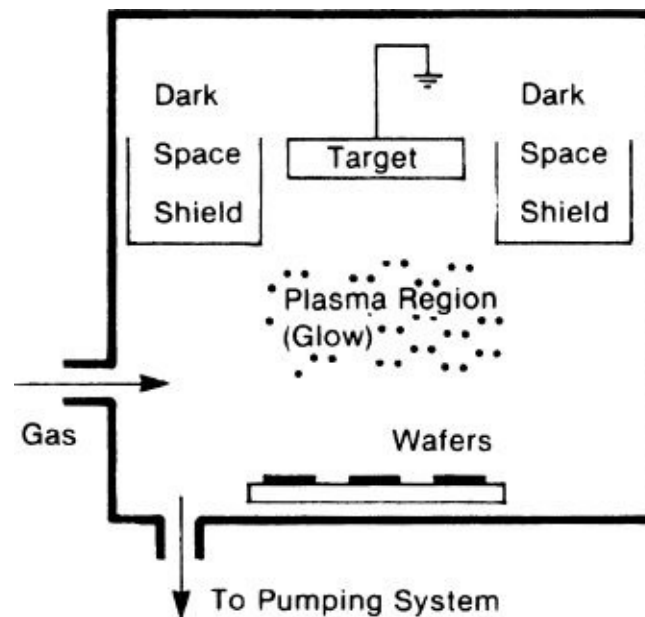


FIGURE 13.10 Typical sputtering equipment.

There are several advantages of sputtering over vacuum evaporation. One is the aforementioned conservation of target material composition. A direct benefit of this feature is the deposition of alloys and dielectrics. A 2 percent aluminum or copper target material yields an aluminum and a 2 percent copper film on the wafers.

Step coverage is also improved with sputtering ([Fig. 13.11](#)). Whereas, evaporation proceeds from a point source, sputtering is a planar source. Material is being sputtered from every point on the target, with material arriving at the wafer

holder with a wide range of angles to coat the wafer surface. Step coverage is further improved by rotating the wafer holder and by heating the wafer.

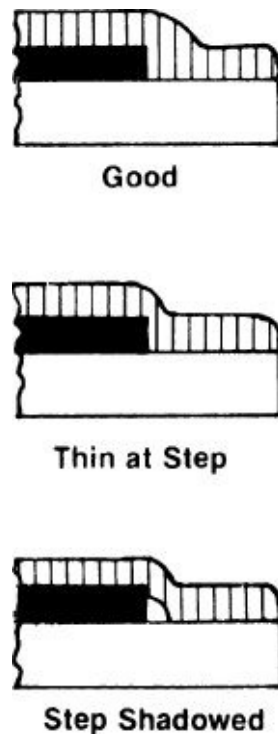


FIGURE 13.11 Step coverage.

Adhesion of the sputtered film to the wafer surface is also improved over evaporation processes. The higher energy of the arriving atoms makes for a better adhesion, and the plasma environment inside the chamber has a “scrubbing” action of the wafer surface that enhances adhesion. Adhesion and surface cleanliness can be increased by grounding the wafer holder and sputtering the wafer surface for a brief time prior to the deposition. In this mode, the sputter system is functioning as an ion-etch (sputter-etch, reverse-sputter) machine, as described in [Chap. 10](#).

Another technique to improve step coverage and uniform film formation in deep holes is a collimated beam ([Fig. 13.12](#)). Atoms come off of the target at many angles and tend to fill the sides of holes before filling the bottom. A collimator is a physical barrier plate similar to a honeycomb with round or hexagonal holes. It is grounded for electrical neutrality. Atoms arriving at the collimator at steep angles are caught on the sides, while straighter-angle atoms continue onto the wafer surface. The thickness of the collimator is a factor in the degree of collimation of the atom beams.

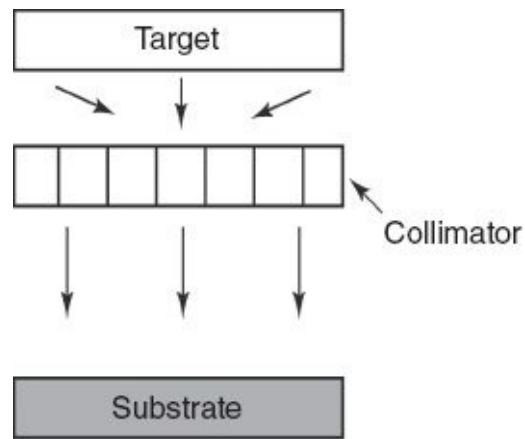


FIGURE 13.12 Sputtering with a collimator.

Uniform film coverage in deep high aspect holes is always achieved with a collimator system. Normally, sputtered target materials are atoms. Researchers discovered that introducing metals into the plasma created ions. Also placing a bias on the wafers attracted the metal ions directly into the holes, providing more uniform coverage. The process is called ionized deposition or I-PVD. Moreover, there is a secondary sputtering (resputter) occurring at the bottom of the hole. First a metal layer is laid down and incoming ions effectively sputter the bottom layer that in turn deposit onto the side of the hole^u ([Fig. 13.13](#)).

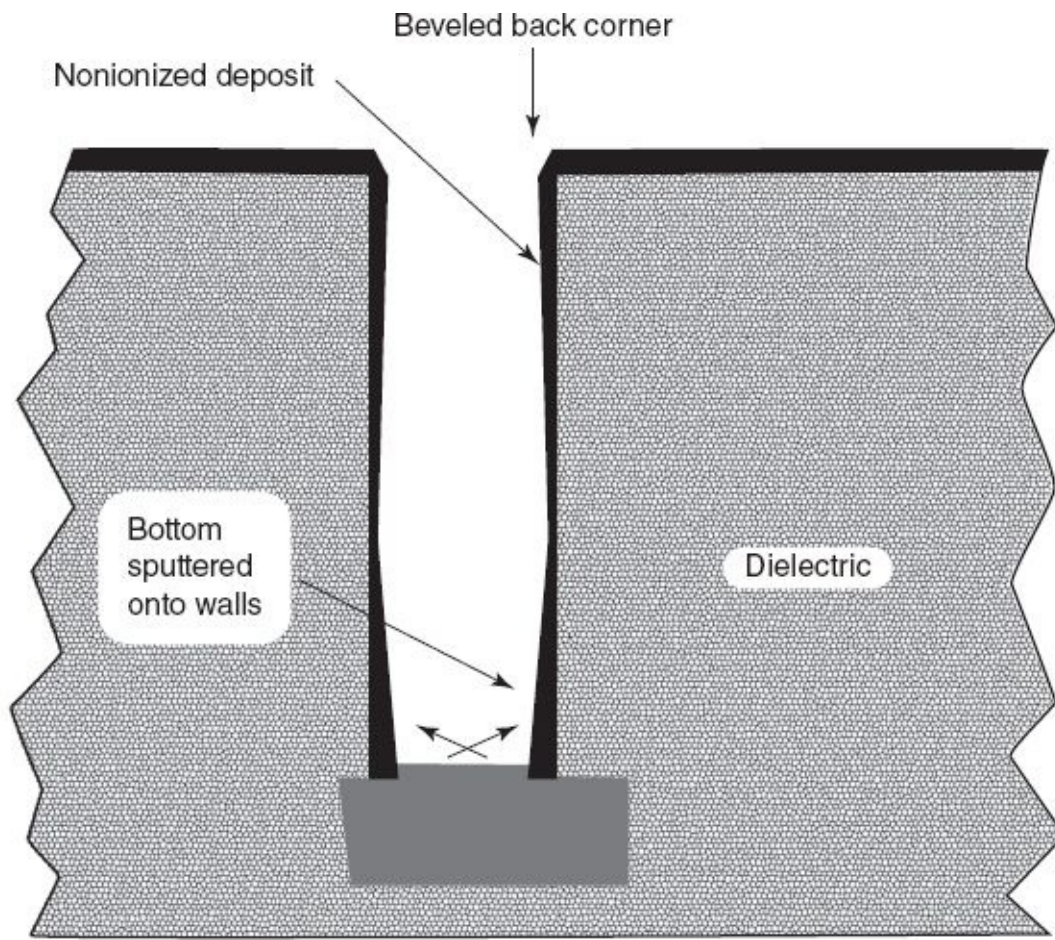


FIGURE 13.13 Ionized PVD showing the effects of resputtering.

Perhaps the greatest contribution of sputtering is the control of film characteristics available by the balancing of the sputtering parameters of pressure, deposition rate, and target material. Sandwiches of material can be sputtered in one process with multiple target arrangements.

Clean and dry argon (or neon) is required to maintain film composition characteristics, and low moisture is required to prevent unwanted oxidation of the deposited film. The chamber is loaded with the wafers, and the pressure is reduced by pumps (pumped down) to the 1×10^{-9} -torr range. The argon is introduced and ionized. Control of the argon amount entering the chamber is critical due to its effect of raising the pressure in the chamber. With the argon and sputtered material in the chamber, the pressure rises to a level of about 10^{-3} torr. Chamber pressure is a critical parameter in the deposition rate of the system. After liberating material from the target, the argon ions, the sputtered material, gas atoms, and electrons generated by the sputtering process form a plasma region in front of the target. The plasma region is evident by its purple glow. The plasma region is separated from the target by a darkened region, known as the *dark space*.

There are four sputtering methods available, diode [direct current (dc)], diode [radio frequency (RF)], triode, and magnetron. *Magnetron* sputtering has emerged

as the system of choice. This system uses magnets behind and around the target (Fig. 13.14). The magnets capture and/or confine the electrons to the front of the target and thus to the wafer. Additionally, it minimizes the amount of chamber material that can be sputtered and end up contaminating the deposited film. Magnetron systems are more efficient for increased deposition rates. The resulting ion current (density of ionized argon atoms hitting the target) is increased by an order of magnitude over conventional diode sputtering systems. Another effect is a lower pressure required in the chamber, which contributes to a cleaner deposited film. Magnetron sputtering leaves a lower target temperature, which makes it a favorite for the sputtering of aluminum and aluminum alloys.

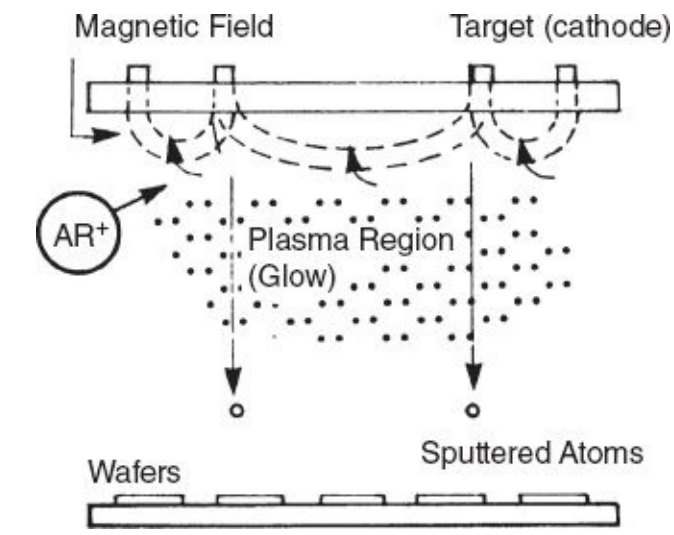


FIGURE 13.14 Magnetron sputtering.

Production-level sputtering systems come in a variety of designs. Chambers are either batch systems or single-wafer in-line designs. Most production machines have load-lock capabilities. A load lock is an antechamber where a partial vacuum is created so that the deposition chamber can be maintained at a vacuum. The advantage of a load lock is a higher production rate. Production machines are usually dedicated to one or two target materials, while development machines have a wider range of capability.

The sputtering process can also accomplish etching and cleaning at the wafer surface. They are achieved by putting the wafer holder at a different field potential from that of the argon, causing the argon atoms to impinge directly on the wafer. This procedure is called *sputter etch*, *reverse sputter*, or *ion milling*. The process removes contamination and a small layer from the wafer. Removal of contamination improves electrical contact between the exposed wafer regions and the film and improves adhesion of the film to the rest of the wafer surface.

Copper Dual-Damascene Process

In the 1990s, IBM introduced the copper-based damascene process to replace aluminum metallization.¹² One of the attractions of copper metallization is that copper can be the plug material, creating a monometal system that minimizes intermetal resistances.

Aluminum metallization ran into a performance barrier as circuits reached 100-MHz speeds. Signals must move fast enough through the metal system to prevent processing delays. Also, longer leads and smaller cross-sections required for larger chips increase the resistance of the metal wiring system. As the number of contact holes increases, the small contact resistance between aluminum and silicon surfaces adds up to become significant. While aluminum provides a workable resistance, it is difficult to deposit in via holes with aspect ratios up to 10:1. To date, barrier metal schemes, stacks, and refractory metals have been employed to reduce aluminum metal system resistance. Additional resistance reductions needed for 0.25- $\mu\Omega$ -cm (and smaller) devices have renewed interest in copper as a conducting metal. Copper is a better conductor than aluminum, with a resistance of 1.7 $\mu\Omega$ -cm, compared to a 3.1 $\mu\Omega$ -cm value for aluminum. Copper is resistant to electromigration and can be deposited at low temperatures. It also can be used as a plug material. Deposition can be by CVD, sputtering, electroless plating, and electrolytic plating. Drawbacks, besides lack of a learning curve, include etching problems, vulnerability to scratching, corrosion, and the requirement of barrier metals to keep the copper out of the silicon. Nevertheless, the overall benefits of copper, led IBM, followed quickly by Motorola, to announce the availability of production copper-based devices in 1998.¹³ Current integrated circuits (ICs) are being developed with copper metallization and low-k dielectrics. The primary benefits are increased performance and a reduction in the number of metal layers required.

Low-k Dielectric Materials

In the dual-damascene illustration, the dielectric separating the two metals is silicon dioxide. However, this material presents a problem for high-performance circuits. The slowing of circuit signals results from the combination of the metal resistance (R) and capacitance (C). It is called the *RC constant* of the system. A major contribution to the capacitance factor is the dielectric constant of the material used to separate the metal layers, the intermediate metal dielectric (IMD).

Silicon dioxide has a dielectric constant (k) in the 3.9 range. Successful circuits will require k values dropping to the 1.5 to 2.0 range, according the SIA *International Technology Roadmap for Semiconductors*. In addition to the dielectric property the IMD must have a number of chemical and mechanical properties. They

include thermal stability (subsequent metal processes can take an initial film through a number of heat steps up to the 450°C), good etch selectivity, being pinhole free, enough flexibility to withstand on chip stresses, and compatibility with the other processes.

A number of low-k dielectric materials have been developed to meet ULSI circuit needs. They are listed in [Fig. 13.15](#). The major categories are oxide-based materials, organic-based, and variations of each. The organics, based on poly(arylene)ethers (PAEs) or hydrido-organic siloxane polymers (HOSPs), offer the advantage of spin-on applications. Spin-on processes can provide great uniformity and planarity, and they are less expensive than CVD processes.

Metal system	Low-k material
Aluminum	arylene HSQ methyl silsesquioxane F-doped oxide F-doped amorphous carbon Arylene-F (AF ₄) Xerogel
Gold	polyimide BCB Xerogel

FIGURE 13.15 Low-k materials. (Source: Future Fab International.)

The Dual-Damascene Copper Process

Transitioning from aluminum to copper metallization is not a simple matter of switching materials. Copper has its own set of problems and challenges. It is not easily etched by wet or dry techniques. Copper has a high electrical resistance connection with silicon. It diffuses easily through silicon dioxide and can enter the silicon structure. There, it can degrade device performance and create junction leakage problems. Copper does not adhere well to silicon dioxide surfaces, causing structural problems. These challenges led to the development of a unique process specifically designed to overcome the copper problems and produce a high-production process. It features a lithography process, the development of a low-k barrier or liner process, copper electrochemical plating, and a chemical-mechanical polishing process.

The damascene process was introduced in [Chap. 10](#). But its origins go back to the Middle Ages and a metal inlay process used for decorating pottery. It was developed

around the ancient city of Damascus, hence the damascene name. The concept is simple. A trench is formed in a surface dielectric layer using a photolithography process, and the required metal is deposited into it. Usually, the trench is overfilled, requiring a CMP step to replanarize the surface ([Fig. 13.16](#)). This process offers superior dimensional control, because the trench width defines the metal width. It eliminates the width variation from a typical metal-etch process following a metal deposition.

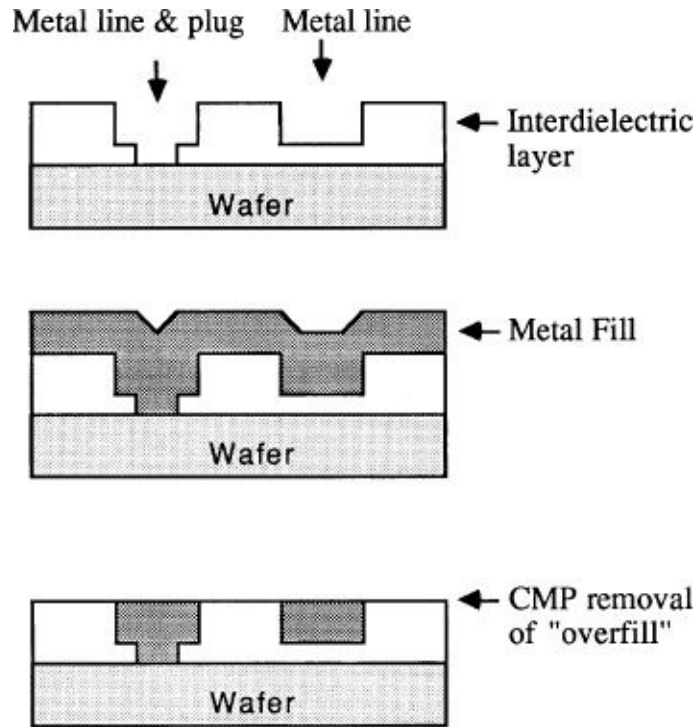


FIGURE 13.16 Dual-damascene (inlaid) process.

In practice the process is a bit more complicated. In a multilayer metal system, there has to be a direct electrical connection from the first metal layer to the device. And there has to be a second patterning to create a trench to carry the second-layer metal, hence, the name *dual-damascene*. [Figure 13.17](#) illustrates a typical dual-damascene process that creates two metal levels. It starts with the first metal already in place. A layer of a low-k dielectric is deposited and planarized with a CMP process. A patterning step creates a via hole in the dielectric layer and a “trench” in the dielectric. A second patterning step results in the lowering of the dielectric and a “step back” on the surface to allow a wider trench width. This pattern leaves a wider opening in the top layer to allow enough width for the copper strip to carry the required current levels. This sequence offers the advantage of a one-step process to fill the via and form the copper metal lead. There are a number of variations on this basic dual-damascene process, each ending up with a narrow via and a wider trench

opening ready for metal fill.

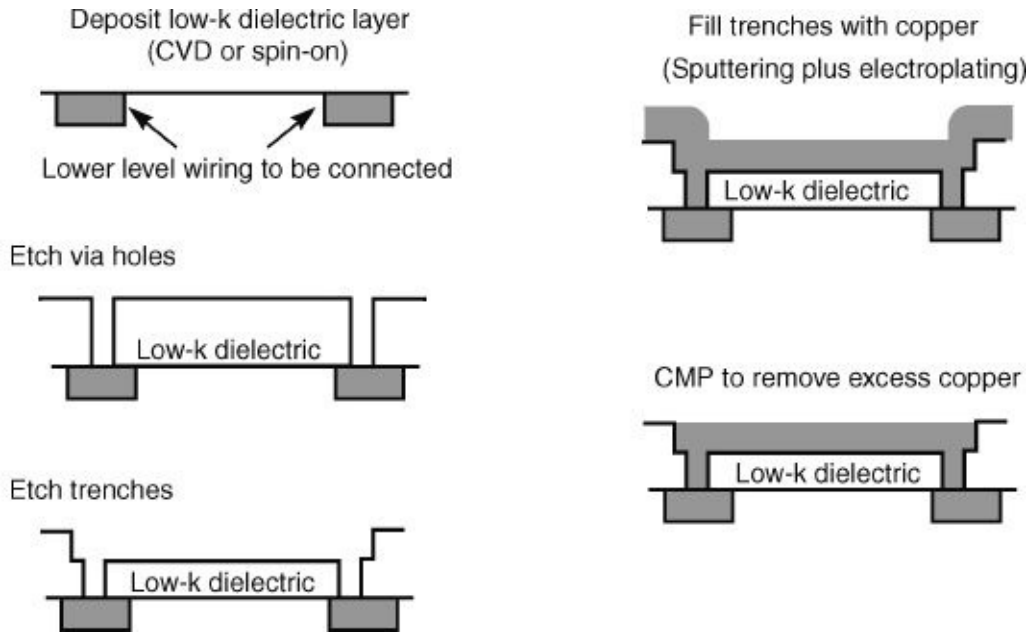


FIGURE 13.17 Typical dual-damascene process.

Barrier or Liner Deposition

As previously mentioned, copper diffuses easily through silicon dioxide and can cause electrical performance problems if it gets in the circuit components. This problem is addressed by depositing a “liner” layer in the via hole bottom and sides ([Fig. 13.18](#)). Typically, the material is tantalum (Ta), 50-to 300-Å thick.¹⁴ Depending on the material, either sputtering or CVD deposition is used to create the barrier or liner. These vias have very high aspects and challenge the process to produce a uniform film over the entire inside surface of the via and trench.

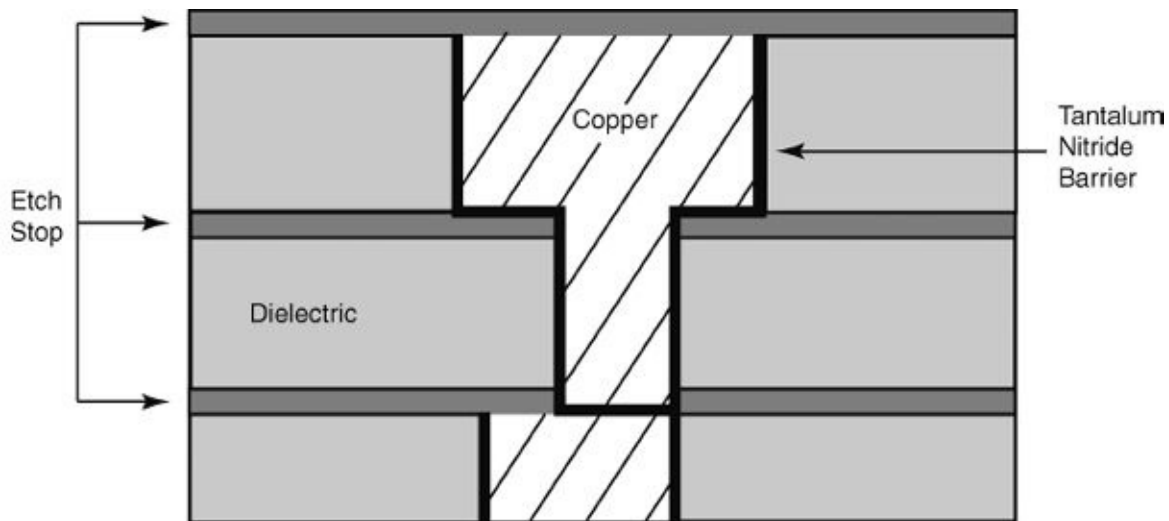


FIGURE 13.18 Single-level dual-damascene with tantalum nitride barrier. (From Wolf and Tauber, *Microchip Manufacturing*, Lattice Press.)

Seed Deposition

Copper can be deposited by sputtering, or CVD deposition, but electrochemical plating (ECP) has emerged as the preferred deposition method. Producing a uniform, void-free copper film with ECP requires a starting “seed” layer in the via or trench hole. PVD techniques are used to deposit the copper seed (300 to 2000 Å)¹⁵ in the via hole. The challenge, as in the barrier or liner deposition, is producing a uniform layer in very high-aspect vias.

Electrochemical Plating

Electroplating has emerged as a production copper deposition method due to its low temperature and low cost.¹⁶ Low temperature is necessary when low-k dielectric layers are used. The seed layer must uniformly coat the bottom and sides of the via/trench to ensure uniform physical and electrical properties of the copper metal lead. Electroplating of copper has been a mainstay of printed circuit board processing for decades ([Fig. 13.19](#)). The wafer is suspended in a bath containing copper sulfate (CuSO_4) and is connected to a cathode (negative pole). With the application of a current, the bath components disassociate. The copper “plates-out” (deposits) on the wafer and hydrogen gas is liberated at the anode. One concern is the uniformity of the film across the wafer. The varied materials and structures on the wafer surface mitigate against uniform current distribution. Nonuniform film growth and density can be the result. Another concern is buildup of a bevel at the lip of the opening. This is addressed by a separate cleaning step after the plating. Nonuniform areas across the wafer surface will have different removal rates in the CMP process. Production level ECP systems will include wafer preclean, the plating sections, bevel removal, and annealing.

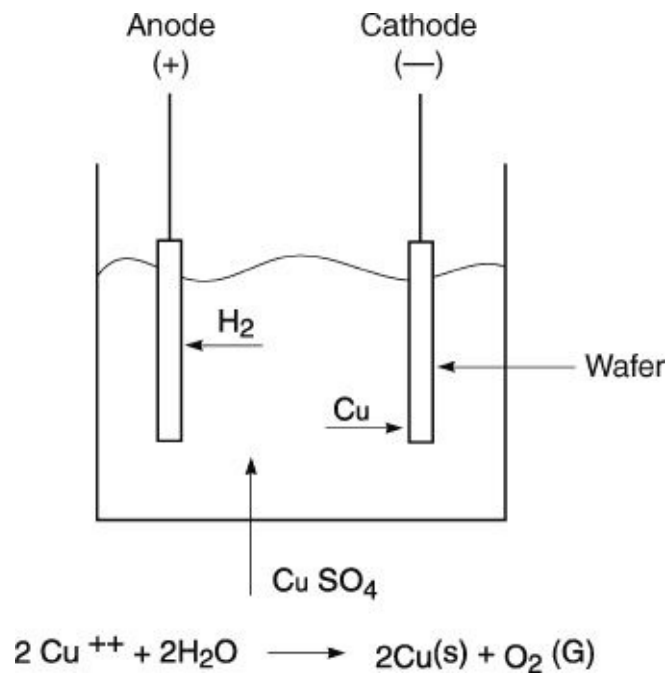


FIGURE 13.19 Schematic of electroplating of copper.

Chemical-Mechanical Processing

Chemical-mechanical processing is used in several steps in the semiconductor process. In [Chap. 3](#), its use for planarizing raw silicon wafers was described. In [Chap. 10](#), we described its use for planarizing in-process wafers to achieve a flat surface for lithography accuracy. The post-copper CMP is a similar process but with a different surface to be flattened. During copper plating, the via or trench hole is overfilled to ensure complete filling of the trench. Before proceeding to the next step, it is necessary to reflaten the surface by removing the copper overfill. The process details are discussed in [Chap. 10](#).

CVD Metal Deposition

Doped Polysilicon

The advent of silicon-gate MOS technology resulted in deposited polysilicon lines on the chip used as conductors. For use as a conductor, the polysilicon has to be doped to increase its conductivity. Generally, the preferred dopant is phosphorus, due to its high solid solubility in silicon. Doping is by either diffusion, ion implantation, or *in situ* doping during an LPCVD process. Each of the methods produces a different doping result. The differences relate to the doping temperature's effect on the grain structure. The lower the temperature, the greater the amount of dopant trapped in the polygrain structure, where it is unavailable for

conduction. This is the situation with ion implantation. Diffusion doping results in the lowest film-sheet resistivity. *In situ* CVD doping has the lowest dopant carrier mobility due to grain-boundary trapping.

Doped polysilicon has the advantage of a good ohmic contact with the wafer silicon and can be oxidized to form an insulating layer. Polysilicon oxides are of a lower quality than thermal oxides grown on a single-crystal silicon because of the nonuniformity of the oxide grown on the rougher polysilicon surface.

Although polysilicon has a low-contact resistance with silicon, it still exhibits too high a resistance to the other metal material(s) used for conductive leads. The problem is addressed by creating a multimetal stack of the polysilicon and a silicide (such as titanium silicide). These are called polycide structures (see [Chap. 16](#)).

CVD Refractory Deposition

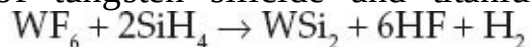
Advances in LPCVD offer a third choice for metal depositions. LPCVD offers the advantages of not requiring expensive and maintenance-intensive high-vacuum pumps, has good conformal step coverage, and high production rates. Perhaps the most-often deposited CVD refractory metal film is tungsten (W).

Tungsten is used in a variety of structures, including contact barriers, MOS-gate interconnects, and via plugs. The filling of via holes is a key to effective multimetal systems. The dielectric layer is relatively thick and the via holes have to be relatively narrow (high-aspect ratio). These two factors make for a difficult continuous metal deposition to fill the vias without thinning of the metal in the via. Selective CVD-deposited tungsten plugs fill the entire via and create a planar surface for a subsequent conducting metal layer. For use as a barrier metal, tungsten can be deposited selectively by the silicon reduction of the gas tungsten hexafluoride (WF_6) by the following reaction: $2\text{WF}_6 + 3\text{Si} \rightarrow 2\text{W} + 3\text{SiF}_4$

Tungsten can also be deposited selectively over aluminum and other materials from WF_6 . The processes are called *substrate reduction*. Tungsten is also deposited from WF_6 by hydrogen reduction; the reaction is: $\text{WF}_6 + 3\text{H}_2 \rightarrow \text{W} + 6\text{HF}$

All of the depositions are performed in LPCVD systems at temperatures of about 300°C , which makes the processes compatible with aluminum metallization.

The depositions of tungsten silicide and titanium silicide proceed by the



following reactions: $\text{TiCl}_4 + 2\text{SiH}_4 \rightarrow \text{TiSi}_2 + 4\text{HCl} + 2\text{H}_2$

Metal-Film Uses

MOS Gate and Capacitor Electrodes

Most electrical devices depend on the passage of an electrical current to operate. Capacitors are an exception. These devices (see [Chap. 16](#)) require two conductive layers, called electrodes, separated by a dielectric. In most designs, the top electrode is a section of the conductor metal system. A discussion of the relationship of capacitor parameters is in [Chap. 2](#).

MOS transistors are a capacitor structure, and the top electrode, called a *gate*, is a critical structure in MOS circuits.

Backside Metallization

A metal layer is sometimes sputter-deposited onto the entire back of the wafer for the packaging process. The metal functions as a thermal interconnection layer or bonding in certain packaging processes. An array of metals are used including gold, platinum, titanium, and copper (see [Chap. 18](#)).

Vacuum Systems

At the dawn of microchip manufacturing, there were only two vacuum-based processes: the evaporation of aluminum and backside gold. Today about a quarter of the processes are done in vacuum or low pressure. They include lithography exposure, strip and etching systems, ion implantation, sputtering processes, LPCVD, PECVD, and rapid thermal processing. Additionally, automated processing requires low-pressure environments for load-lock stations and transfer tools. Vacuum chambers provide process conditions free of contaminating gases. In the deposition processes, the vacuum increases the mean free path of the depositing atoms and molecules, which in turn results in more uniform and controllable deposited films. LPCVD takes place in the pressure range down to 10^{-3} torr (medium range), while the other processes take place at pressure ranges down to 10^{-9} torr (high to ultra-high range). Medium range is reached with mechanical vacuum pumps. These same pumps are used to initially reduce the pressure in the high-vacuum process chambers. In this role, they are called *roughing pumps*. Additionally, mechanical vacuum pumps are used on the outlet end of high-vacuum pumping systems to assist in the removal of gas molecules from the pump to the exhaust system.

After the rough vacuum is established, a high-vacuum pump takes over to establish the final vacuum. The industry has gone through a succession of hi-vac pumps (*oil diffusion*, *cryogenic*, *ion*) finally settling on *turbomolecular*. All pumps are constructed of materials that will not *outgas* into the system and compromise the vacuum. Materials used are typically type 304 stainless steel, oxygen-free high-conductivity copper (OFHC), Kovar, nickel, titanium, borosilicate glasses,

ceramics, tungsten, gold, and some low-vapor-pressure elastomers. Pumps used to evacuate corrosive and toxic gases or reaction byproducts must have corrosion-free inside surfaces. Also, care must be taken in servicing pumps with these types of applications.

Pumps are selected and used based on a number of criteria, including: • Vacuum range required • Gases to be pumped (lighter gases such as hydrogen are more difficult to pump) • Pumping speed

- Overall throughput
- Ability to handle impulsive loads (periodic outgassing) • Ability to pump corrosive gases • Service and maintenance requirements • Downtime
- Cost

Recall from [Chap. 2](#) that pressure in a system results from the activity of gas atoms or molecules in an enclosure striking the chamber walls with some force. Reduction of the pressure in a system requires the removal of the gas in the chamber. This is generally accomplished by the pump establishing a lower pressure, first within itself, which allows gas material in the process chamber to flow into the pump, where it is removed entirely from the system. At very low pressures, there is not much material in the chamber, and continued pressure reduction requires that the system be leak-free and not add to the pressure by its own outgassing. Some systems use cold traps to prevent material from the pump from *backstreaming* into the chamber.

Dry Mechanical Pumps

Replacing the earlier oil mechanical pumps are *dry mechanical pumps*. Oil-based pumps are a source of contamination as the oil absorbs exhausted gases. The toxic gases represent particular safety problems. Dry pumps are based on a “roots” design. These are screw or claw designs that mechanically “grab” the gases reducing pressure in the process chamber before the hi-vac pump takes over.

Turbomolecular Hi-Vac Pumps

Turbomolecular pumps are similar in design to a jet turbine engine. A series of blades ([Fig. 13.20](#)) with openings are mounted and rotated at very high speeds (24,000 to 36,000 rpm)¹⁷ on a central shaft. Gas molecules from the chamber encounter the first blade and gain momentum from the collision with the rotating blade. The momentum direction is downward to the next blade, where the same thing happens. The net result is a removal of gas from the chamber. The use of a momentum transfer makes the pumping principle the same as an oil-diffusion pump.

Major advantages of turbomolecular pumps are a lack of backstreaming from oils, no need to recharge, high reliability, and pressure reduction into the high vacuum range. Drawbacks are a slower pumping speed compared to oil diffusion and cryogenic pumps and vibration and wear due to the high rotational speeds. An addition to turbo pumps is a *drag-type* pump. Molecules are bounced off a rotating drum or disk rather than vanes or stators.¹⁸

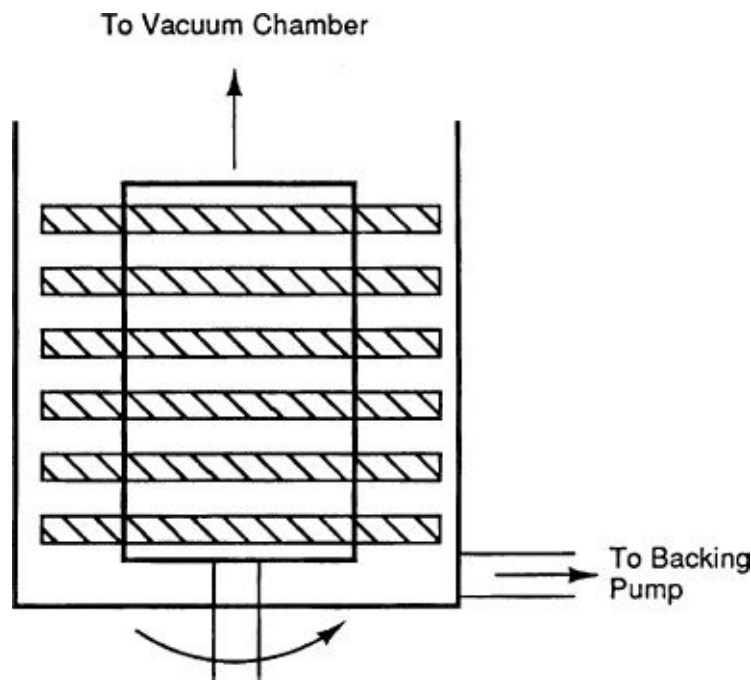


FIGURE 13.20 Turbomolecular pump.

These *combination* pumps can exhaust at high pressures. Use of turbo pumps with corrosive gas processes requires coating the rotors and stators and/or heating the pump to keep the gases from forming solid particles that deposit on the pump parts.

Review Topics

Upon completion of this chapter, you should be able to: 1. List the requirements of a material for use as a chip-surface conductor.

2. Draw cross-sections of single-and two-layer metal schemes.
3. Describe the purpose and operation a low-k dielectric layer.
4. Make a list of three materials used in the metallization of semiconductor devices and identify their specific use(s).
5. Describe the principle of sputtering.
6. Draw and identify the parts of a sputtering system.
7. Describe the principle and operation of turbo and cryogenic high-vacuum

pumps.

References

1. *International Technology Roadmap for Semiconductors*, 2005 Executive Summary:79.
2. Wolf, S. and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Sunset Beach, CA:332.
3. Riley, P., Peng, S., and Fang, L., "Plasma Etching of Aluminum for ULSI Circuits," *Solid State Technology*, PennWell Publishing, Feb. 1993:47.
4. Sze, S. M., *VLSI Technology*, McGraw-Hill, 1983, New York, NY:347.
5. Singer, P., "New Interconnect Materials: Chasing the Promise of Faster Chips," *Semiconductor International*, Nov. 1994:53.
6. Singer, P., "Copper Goes Mainstream: Low-k to Follow," *Semiconductor International*, Nov. 1997:67.
7. Singer, P., "New Interconnect Materials: Chasing the Promise of Faster Chips," *Semiconductor International*, Nov. 1994:54.
8. Brown, D. M., "CMOS Contacts and Interconnects," *Semiconductor International*, 1988:110.
9. Tisdale, G., "Next-Generation Aluminum Vacuum Systems," *Solid State Technology*, May 1998:79.
10. Pramanik, D. and Jain, V., "Barrier Metals for ULSI," *Solid State Technology*, PennWell Publishing, Jan. 1993:73.
11. Singer, P., "Copper Goes Mainstream: Low-k to Follow," *Semiconductor International*, Nov. 1997:68.
12. Braun, A., "ECP Technology," *Semiconductor Technology*, May 2000:60.
13. Pauleau, Y., "Interconnect Materials for VLSI Circuits," *Solid State Technology*, Feb. 1987:61.
14. Aronson, A. J., "Fundamentals of Sputtering," *Microelectronics Manufacturing and Testing*, Jan. 1987:22.
15. Ballingall, J., "State-of-the-Art Vacuum Technology," *Microelectronics Manufacturing and Testing*, Oct. 1987:1.
16. Wolf, S., *Microchip Manufacturing*, 2004, Lattice Press, Sunset Beach, CA:334.
17. Wolf, S. and Tauber, R., *Silicon Processing for the VLSI Era*, 1986, Lattice Press, Sunset Beach, CA:95.
18. Singer, P., "Vacuum Pump Technology Leaps Ahead," *Semiconductor International*, Cahners Publishing, Sep. 1993:53.

CHAPTER 14

Process and Device Evaluation

Introduction

The wafer-fabrication process requires a high degree of precision in process control, equipment operation, and process materials manufacture. One process mistake can render the wafer completely useless. One “killer” defect can ruin a die. Throughout the process, a variety of tests and measurements are made to determine both wafer quality and process performance. The tests take place on in-process wafers, test die and production die, and the finished circuit. Individual tests are described in this chapter. Statistical process control basics are addressed in [Chap. 15](#).

Metrology is the general term applied to the measurement of physical surface features. In wafer fabrication these include pattern widths, film depths, defect identification and location, and pattern registration errors. Good characterization can warn of a process that is about to go out of control, and device characterization is essential to analyze circuit performance and conformance to operating specifications. Consequently, every process step has a rigid set of equipment and process parameters that are controlled (temperature, time, and so forth). After every significant process step, there is an evaluation of the result on the wafer or a test wafer. *Test wafers* (or *monitor wafers*) are blank wafers that are included in the process step for postprocess measurements. Many of the tests are destructive and cannot be performed on the device wafers or cannot be performed on the actual components in the chip. In [Chaps. 7–13](#), the important parameters for each process were identified (e.g., film thickness, resistivity, cleanliness). Here, the basic theory, applicability, and range of sensitivity of the test methods are examined.

Some are direct measurements, and some are indirect. One group includes electrical measurements of test wafers and on the actual devices. They measure the direct effect of some of the processes, such as ion implantation. Device performance measurements are usually inclusive of several processes, and the results are used to infer individual process-parameter control. Another group directly measures physical parameters such as layer thicknesses and widths, composition, and others. This group includes defect detection. A third group measures contamination in and on the wafer and in materials.

Not surprisingly, test and measurement methods have changed along with the levels of integration and smaller image sizes. Ultra-large-scale integration (ULSI) technology is ushering in nanometer-and angstrom-level inquiry, called the *nanoanalysis era*.¹ And the price of in-line testing is going up. Larger wafers and more dense circuits require more tests to properly characterize processes. High-volume processing requires real-time testing and analysis to guard against scrapping volumes of production wafers. Data-management systems for ULSI circuits usually include onboard statistical analysis and database management

capabilities.

Wafer Electrical Measurements

Resistance and Resistivity

The object of the fabrication process is to form, in and on the wafer surface, solid-state electrical components (transistors, diodes, capacitors, and resistors) that are wired (surface “wiring”) together to form the circuit. Each of the components must meet individual electrical performance specifications if the entire circuit is to function. Throughout the process, electrical measurements are performed to judge both the process and predict electrical device performance.

Resistivity Measurements

The addition of dopants to the wafer, both during crystal growth and during the doping processes, alters the electrical characteristics of the wafer. The altered parameter is its resistivity, which is a measure of a material's specific "resistance" to the flow of electrons ([Fig. 14.1](#)). Whereas the resistivity of a given material is a constant, the resistance of a specific volume of the same material is a function of its dimensions and resistivity. This relationship parallels that of density and weight. For example, the density of steel is a constant, whereas the weight of a particular piece depends on its volume.

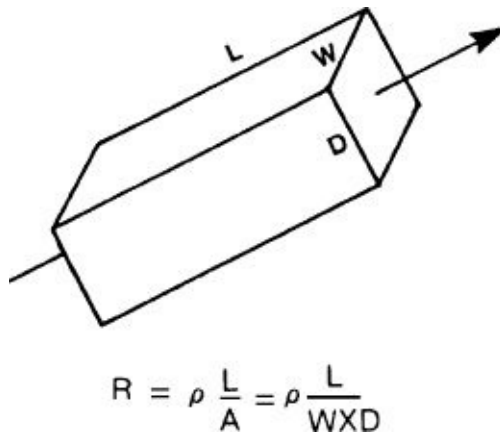


FIGURE 14.1 Relationship of resistance to resistivity and dimensions.

The units of resistance (R) are ohms (Ω) and the units of resistivity (ρ) are ohm-centimeters ($\Omega\text{-cm}$). Because adding dopants to a wafer will alter its resistivity, measurement of resistivity is actually an indirect measure of the amount of dopants added.

Four-Point Probe

The parameters of resistance, voltage, and current are governed by Ohm's law. The three parameters are related mathematically in the following way:

$$R = V/I = (\rho) L/A = (\rho) L/(W \times D)$$

where R = resistance

V = voltage

I = current

ρ = resistivity of sample
 L = length of sample
 A = cross-sectional area of sample
 W = width of sample
 D = depth of sample

Theoretically, the resistivity of a wafer can be measured with a multimeter ([Fig. 14.2](#)) by measuring the voltage at a constant current through a sample of known dimensions and calculating the resistivity. However, the resistance between the probes and the wafer material is too great to accurately measure the resistivity of semiconductors with their relatively low quantity of dopants.

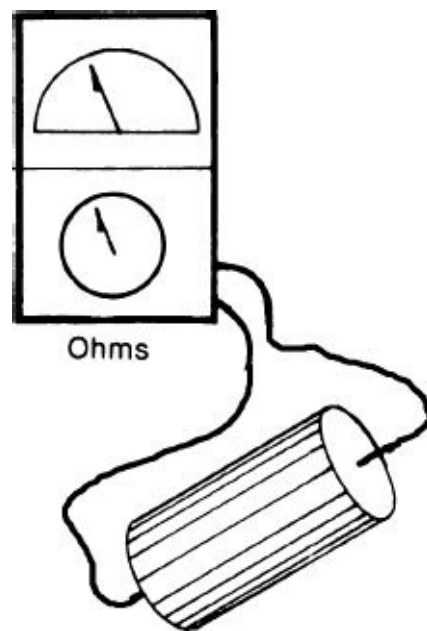


FIGURE 14.2 Multimeter.

The four-point probe is the instrument used to measure resistivity on wafers and crystals. It employs four thin, in-line probes connected to a power supply and voltmeter. The four-point probe consists of four thin metal probes arrayed in a line. The two outside probes are connected to a power supply, and the inside probes are connected to a voltage meter. During operation, the current passes between the outer probes, and the voltage drop (change) is measured between the inner probes ([Fig. 14.3](#)). The relationship of the current and voltage values is dependent on the resistance of the space between the probes and the resistivity of the material. The four-point probe cancels out the effects of probe-wafer contact resistance on the measurement.

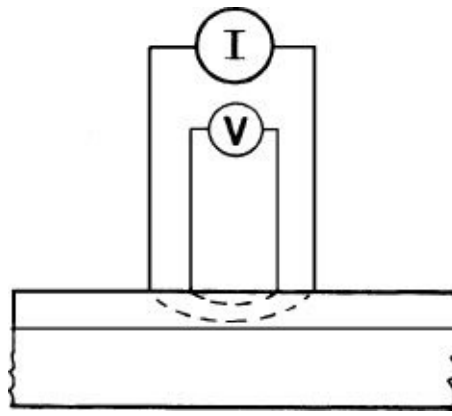


FIGURE 14.3 Four-point probe measurement of a thin layer.

Process and Device Evaluation

Using a four-point probe, the voltage and current are related to the resistivity by the following relationship:

$$\rho = 2\pi sV/I$$

where s is the distance between probes, when s is less than the wafer diameter and less than the film thickness.

$$R_s = 4.53V/I$$

Sheet Resistance

The four-point probe measurement just described is used to measure the resistivity of wafers and crystals. It is also used to measure the resistivity of thin layers of dopants added into the wafer surface by the dopant processes. When a four-point probe measurement is made on a thin layer of added dopants, the current is confined in the layer ([Fig. 14.3](#)). A thin layer is defined as a layer thinner than the probe spacing (distance between probes).

The electrical quantity measured on a thin layer is called *sheet resistance*, R_s . This quantity has the units of ohms per square ($\Omega/\square$). The concept of ohms per square can be understood by considering the resistance of two squares of the same thin material of equal thickness ([Fig. 14.4](#)). Since the resistivity of r is the same for each piece, and $T_1 = T_2$, the sheet resistance is the same for each piece. Or, the resistance of the thin sheet is a constant for any square of the same material.

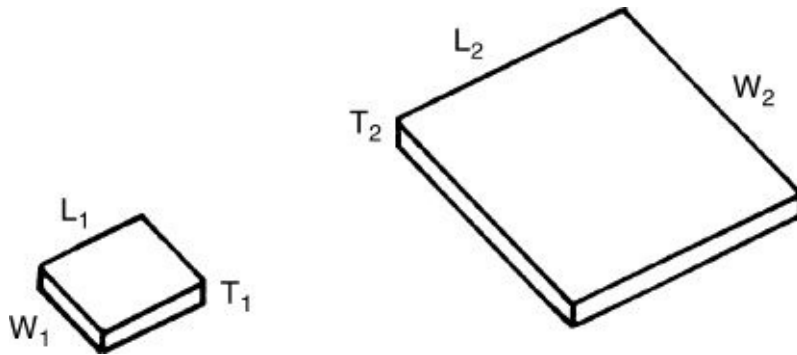


FIGURE 14.4 Resistance of a “square.”

The formula relating sheet resistance to the voltage and current is where 4.53 is a constant that arises from the probe spacing. Some companies elect to drop the constant 4.53 from the formula and just measure the V/I of a wafer as in [Fig. 14.5](#).

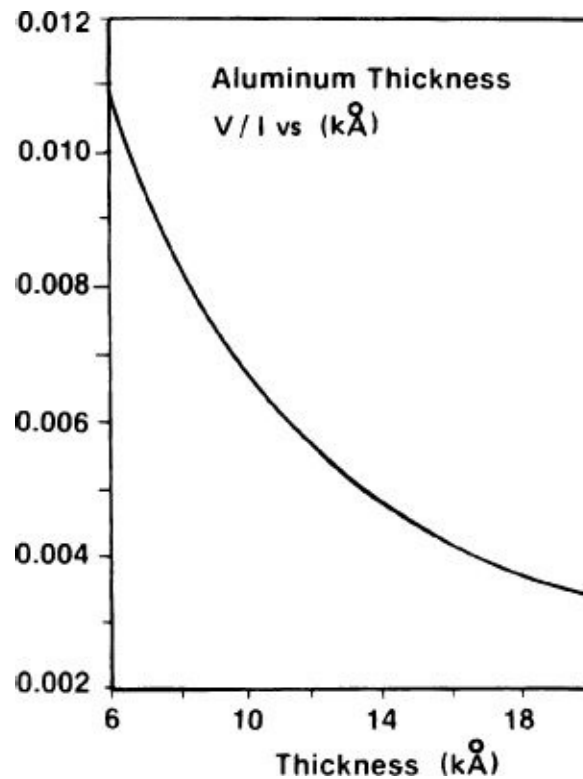


FIGURE 14.5 Voltage/current (V/I) versus thickness of aluminum.

Four-Point Probe Thickness Measurement

The thickness of uniform conducting layers on an insulating layer can be determined using a four-point probe. For thin films, the formula is:

$$T = \rho_s / R$$

where T = layer thickness

ρ_s = resistivity

R_s = sheet resistance

Since the resistivity is a constant for pure materials such as aluminum ([Fig. 14.5](#)), the sheet-resistance measurement is actually a measurement of the film thickness. This formula does not calculate the thickness of a doped layer, since the dopants are not evenly distributed vertically in the doped layer.

Concentration or Depth Profile

The distribution of dopant atoms in the wafer (see [Chap. 11](#)) is a major influence on the electrical operation of a device. The distribution (or dopant concentration profile) is determined by several techniques. One is spreading resistance. After doping, a test-wafer sample is prepared by the bevel technique. After the junction is exposed by the beveling, a series of electrical two-point probe measurements are made sequentially down the bevel ([Fig. 14.6](#)). At each point, the vertical drop of the probes is recorded and a resistance measurement made. The resistance value at each point changes with the change in dopants at each level.

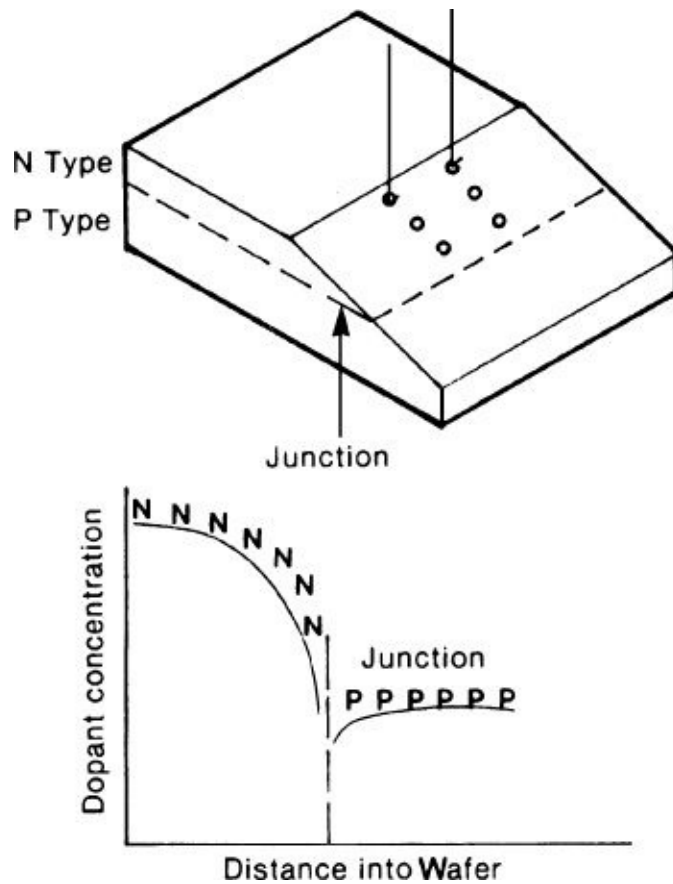


FIGURE 14.6 Spreading resistance.

A computer is used to perform calculations that relate the depth and resistance values to the dopant concentration at each level. The computer uses the data to construct a dopant concentration profile for the sample. This measurement is usually made periodically off-line or when electrical device performance indicates that the dopant distribution may have changed.

Secondary Ion Mass Spectrometry

Secondary ion mass spectrometry (SIMS) is essentially a combination ion-milling and the secondary ion-detection method. Ions are directed at the sample surface, removing a thin layer. Secondary ions are generated from the removed material, which contains the wafers material and the dopant atoms. The ions are collected and analyzed providing a calculation of the amount of dopant at each level, which in turn allows the construction of a dopant profile.²

Optically Modulated Optical Reflection (Thermawave)

Thermawave is the common name of this technique for measuring ion low-dose implant doping amounts. An argon laser beam (1- μm diameter spot) heats a small area of the wafer surface causing an increase in the silicon volume. The volume change in turn changes the optical properties of the surface that is detected by a second laser. Optical property changes come from the amount and species of the dopant implanted.³

Physical Measurement Methods

Product reliability and maintenance of yields requires on-line detection of defects, mistakes, and out-of-control processes to remove out-of-spec wafers from the production line. Process control requires measuring the process results at each of the process steps and knowing the quantity, density, location, and nature of the various problems. This data comes from a series of measurements and evaluations that vary with the sophistication of the circuit in terms of image size, sensitivity to contamination, and density.

Tests are performed on test wafers or directly on product wafers. Product-wafer tests may require revealing the subsurface structure by beveling, microsectioning, or using a focused ion beam (FIB) to remove portions of the circuit.

Layer Thickness Measurements

Color

Both silicon dioxide and silicon nitride layers exhibit different colors on the wafer. We know that, while silicon dioxide is transparent (glass is silicon dioxide), an oxidized wafer has a color. The color seen is actually the result of an interference phenomenon—the same phenomenon that creates the colors of rainbows.

The silicon dioxide layer on a silicon wafer is actually a thin transparent film on a reflecting substrate. Some of the light rays impinging on the wafer surface reflect off the oxide surface, while others pass through the transparent oxide and reflect off the mirrored wafer surface ([Fig. 14.7](#)). When the light rays exit the film, they combine with the surface-reflected ray, giving the surface an appearance of having a color. This phenomenon is the reason oxidized wafers change color as the angle of viewing is changed.

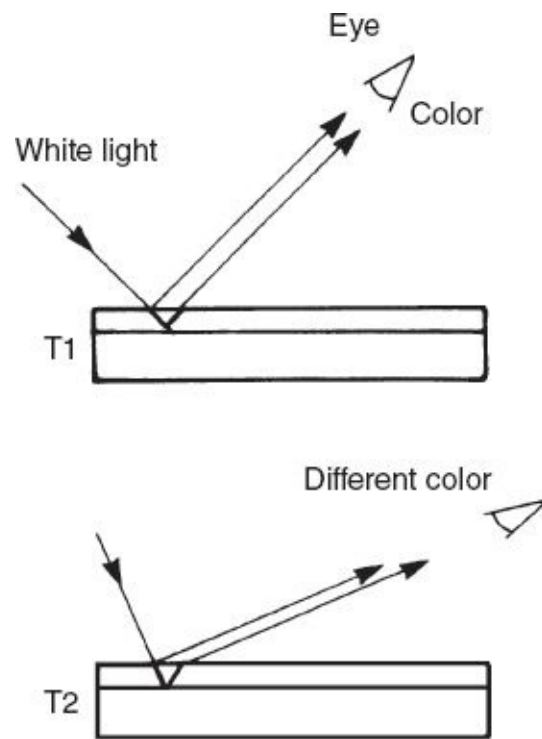


FIGURE 14.7 White-light interference.

The exact color is a function of three factors. One, which is a property of the transparent film material, is the *index of refraction*. The second factor is the *viewing angle*. The third factor is the *thickness of the film*. The color of a thin transparent film becomes an indication of the thickness when the nature of the viewing light is specified (i.e., daylight, fluorescent), along with the viewing angle. The classic color versus thickness chart ([Fig. 14.8](#)) is a regular reference at oxidation and

diffusion stations. Color alone is not an exact indication of thickness because of the consequences of the interference phenomenon.

Film thickness, μm	Color* and Comments
Order I	
0.050	Ten
0.075	Brown
0.100	Dark Violet to Red Violet
0.125	Royal Blue
0.150	Light Blue to Metallic Blue
0.175	Metallic to Very Light Yellow
0.200	Light Gold or Yellow-Slightly Metallic
0.225	Gold with Slight Yellow Orange
0.250	Orange to Melon
0.275	Red Violet
0.300	Blue to Violet Blue
0.310	Blue
0.325	Blue to Blue Green
0.345	Light Green
0.50	Green to Yellow Green
Order II	
0.365	Yellow Green
0.375	Green Yellow
0.390	Yellow
0.412	Light Orange
0.426	Carnation Pink
0.443	Violet Red
0.465	Red Violet
0.476	Violet
0.480	Blue Violet
0.493	Blue
Order III	
0.502	Blue Green
0.520	Green (Broad)
*When viewed perpendicularly in daylight fluorescent light	

FIGURE 14.8 Silicon dioxide thickness color chart.

As the film gets thicker, the colors change in a specific sequence and then repeat themselves. Each repetition of the color is called an *order*. To determine the exact film thickness, a knowledge of the color order is necessary. A principal use of color charts is for process control.

Each oxidation or silicon nitride process is set up to produce a specified thickness. Naturally, the thickness will vary from run to run. Operators quickly become sensitive to the wafer color. When a variation occurs, a quick check of the

chart will indicate if the film thickness is out of specification. Rarely is a process so far off that the film thickness is a whole order (same color, different thickness) out of specification. The accuracy of color chart thickness determination is limited to the accurate perception of the colors (what exactly is red-orange?). A typical chart is accurate to $\pm 300 \text{ \AA}$.

Spectrophotometers or Reflectometry

Film thickness interference-or reflectance-measurement techniques can be automated. To understand the method, let us review the interference effects. Light is actually a form of energy. The interference phenomenon can also be described in terms of energy. White light is really a bundle of rays (different colors), each with different energies. When the rays interfere through the transparent film, the result is a ray of one color, one wavelength, and one energy level. It is our eyes that interpret the energy as a color.

In a spectrophotometer, which is an automatic interference instrument, a photocell takes the place of the human eye. Monochromatic light in the ultraviolet range is reflected off the sample and analyzed by the photocell. To ensure accuracy, readings are made under different conditions. The conditions are changed by either using another monochromatic light (to change wavelength) or changing the angle of the wafer to the beam. Spectrophotometers specifically designed for use in semiconductor technology have onboard computers to calculate the film thickness. With visible and ultraviolet (UV) light sources, these machines can measure films down to the 100-Å level.⁴

Index of refraction measurements are also made with these instruments. Spectrophotometers are also used to measure silicon film thickness. Because silicon is opaque to ultraviolet light, an infrared source is used.

Ellipsometers

Ellipsometers are film-thickness instruments that use a laser light source and operate on a different principle from that of a spectrophotometer. The laser light source is polarized. Polarization is the creation of a wave with all the rays traveling in only one plane. Polarization can be imagined by considering looking into the beam of a flashlight. In an ordinary beam, light rays come to your eyes in many planes, like an arrow with many feathers. A polarized beam has all of the light in only one plane, or an arrow with only one feather ([Fig. 14.9](#)).

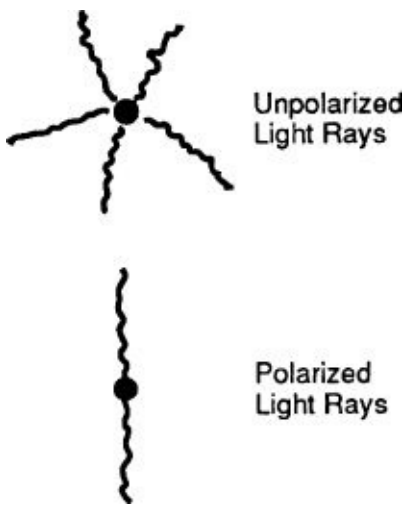


FIGURE 14.9 Unpolarized and polarized light.

In the ellipsometer, the polarized beam is directed to the oxide-covered wafer at an angle. The beam enters the transparent film and reflects off the reflective wafer surface. During its passage through the film, the angle of the beam plane is rotated. The amount of rotation of the beam is a function of the thickness and index of refraction of the film. A detector in the instrument measures the amount of rotation, and an onboard computer calculates the thickness and index of refraction.

Ellipsometers are used to measure thin oxides (50 to 1200 Å) and the index of refraction of the film. Their accuracy in this range is unequaled by any of the other techniques. The technique is also used for multilayer thin-film stacks. Ellipsometer accuracy and range is enhanced with the addition of multiple-viewing-angle capabilities, multiple-wavelength sources, and a reduction in beam spot size.⁵

Oxides thinner than 50 Å can be measured by the corona-oxide-semiconductor (COS) technique described in the electrical measurement section.

Stylus (Surface Profilometers)

Some thin films, such as aluminum, cannot be measured by optical techniques. And in the case of aluminum and other very thin conductive films, the four-point probe thickness measurement is not sufficiently accurate. In these situations, a mechanical moving stylus apparatus is used ([Fig. 14.10](#)). The method requires that a portion of the film be removed, creating a step on a test-wafer surface. This is normally done by a masking and etching step. The prepared sample is mounted and leveled on the pivoting stylus instrument stage.

After leveling, the measuring stylus is lowered gently onto one of the surfaces. The measurement is made as the stage is slowly moved under the stylus. As the stylus goes over the step, its physical position is changed. The stylus is linked to an inductor that generates an electrical signal in response to the stylus vertical position. This signal is amplified and fed into an x-y recorder.

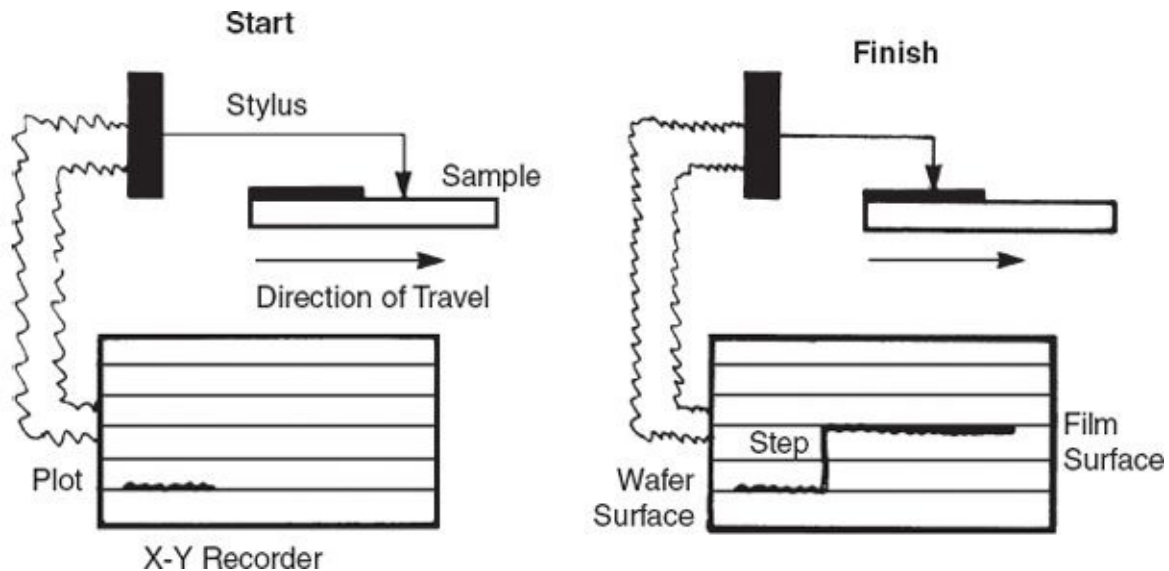


FIGURE 14.10 Step height measurement.

While the leveled wafer is moving under the stylus, it does not move in the vertical direction, and no change in signal is produced. The trace on the x-y plotter is a straight line. When the stylus reaches the surface step, it changes position, causing a change in the signal output. This change is evidenced by a change of pen position on the x-y chart trace. The change in position is relative to the step height, which is read directly from the calibrated x-y chart. Accuracy is related to the tip material and diameter. With diamond tips in the 20-to 50-nm range, surface steps in the nanometer range can be measured ([Fig. 14.11](#)).

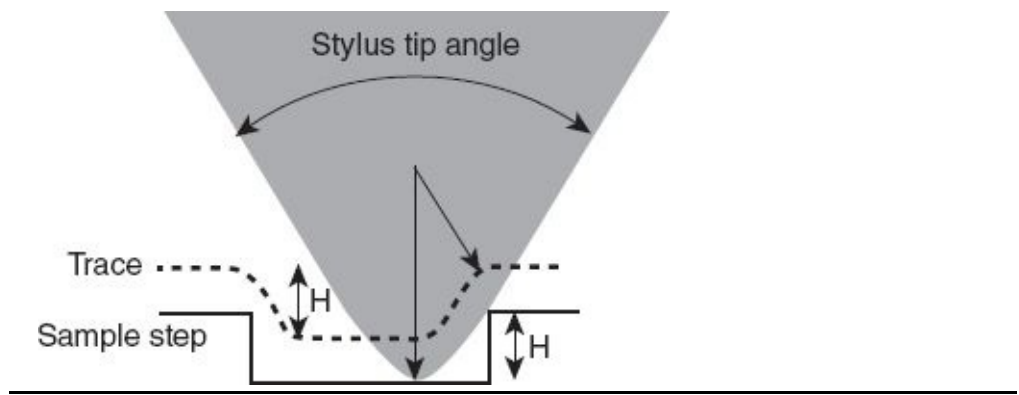


FIGURE 14.11 Stylus profilometers tip or step.

Optical Profilometer

Film thickness or step height is also measured using a noncontact optical profilometer. The setup is similar to a stylus profiler, but a beam of light is swept across the surface and film step and reflected into detectors. The characteristic changes in the reflected beam (or beams) are translated into step height. This technique is explained further in the section on surface profiling.

Photoacoustic

A nondestructive thickness test relies on photoacoustic principles. In 1877, Alexander Graham Bell discovered that, under certain circumstances, the interruption of a light wave will cause a sound. In the semiconductor thickness application, a laser beam is converted to tiny sounds, which are in turn reflected off two surfaces on the wafer surface. By measuring the reflection delay between the two pulses, the thickness can be calculated.

Four-Point Probe

The four-point probe sheet-resistant method can also be used to measure a thin film thickness. The method and calculation is explained in previous sections.

Ultra-Thin MOSFET Gate Thickness

Metal-oxide semiconductor field effect transistor (MOSFET) scaling is taking gate thicknesses down to the 10-nm range. The thickness, area, and material integrity all influence device performance. Measuring all three is critical to process control and device operation. Commercial ellipsometers, with specific laser sources and data processing, can measure silicon dioxide, silicon nitride, and titanium dioxide film thickness in the nanometer range.

Since the levels and interplay of the three variables is critical, and given that these lengths are on the order of atom diameters, *capacitance or voltage* measurements (see the “Diodes” section in this chapter) made on test devices. These plots profile the actual functioning of the device beyond the physical measurement methods.

Gate Oxide Integrity Electrical Measurement

Gate oxide integrity (GOI) encompasses film uniformity and contamination. It is determined by a destructive electrical test called oxide rupture voltage (BV_{ox}). The test structure used is the same as for capacitance-voltage (C/V) analysis. But, in this case, the voltage is continually increased until the oxide is physically destroyed and current flows freely from the gate electrode to the silicon. The maximum voltage that the oxide can withstand before breakdown is a function of its thickness, structural quality, and purity.

Junction Depth

A critical device parameter is the junction depth of the various doped regions. This parameter is measured after each of the doping steps. The measurement methods described are all performed off-line; that is, the test wafers or device wafers have to be taken to a measurement station or laboratory for the measurement.

Groove and Stain

The traditional method of junction depth measurement is by the groove (or bevel) and stain technique. Grooving or beveling is a mechanical method of exposing the junction for viewing and measurement from the horizontal plane ([Fig. 14.12](#)). The extremely shallow depth of the junction requires either grooving or beveling of the wafer to expose the junction.

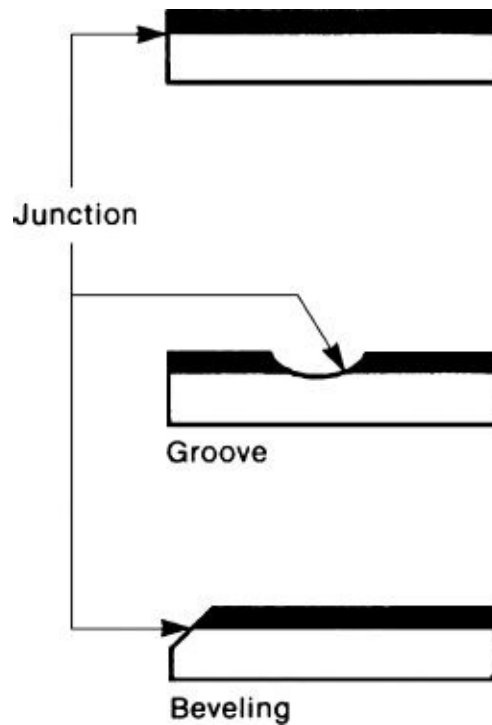


FIGURE 14.12 Exposure of junction by groove or beveling.

The junction itself is not visible to the naked eye. Two techniques, called *junction delineation*, are available to make it visible. Both techniques utilize the electrical differences between N-type and P-type regions. The first technique, the etch technique, starts with the placement of a drop of hydrofluoric acid and water mixture over the junction ([Fig. 14.13](#)). A heat lamp is directed onto the exposed junction. The heat and light cause holes or electrons to flow in each region. As a result of the flowing current, the etch rate of the HF-H₂O mixture is higher on the N-type region, making it appear darker.

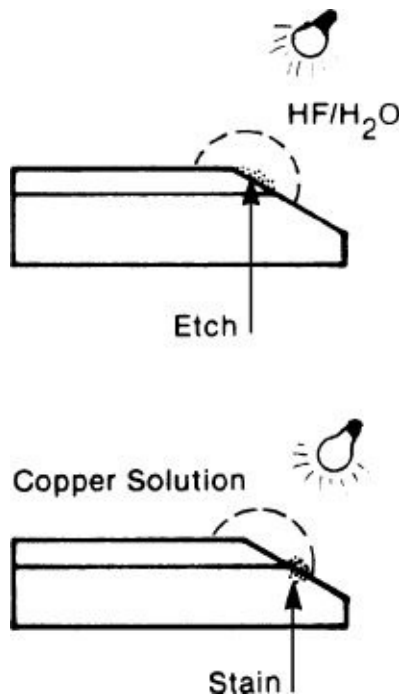


FIGURE 14.13 Etching or staining of junction.

The second delineation technique is electrolytic staining. A mixture containing copper is dropped on the exposed junction. Again, the heat lamp is directed onto the junction. In effect, a battery is formed, with the poles of the junction being the poles of the battery and the copper solution being the electrolytic connection. The current flowing in the liquid drop causes the copper to plate out on the N-type region side of the junction.

The final step, after exposure and delineation of the junction, is depth measurement. A number of methods may be used, including optical interference and the scanning electron microscope (SEM).

Scanning Electron Microscope Thickness Measurement

The SEM technique (described in the following sections) is also used to measure junction depths and film thickness. The wafer is cleaved with the break position over the junction. The exposed wafer junction is delineated by one of the methods described.

The exposed cross-section is positioned to the SEM beam at right angles to the wafer surface and a photograph is taken. The depth is determined from the photograph and the scale factor of the SEM ([Fig. 14.14](#)). SEM and the groove and stain methods provide a visual look at the junction area and lateral diffusion that other methods do not.

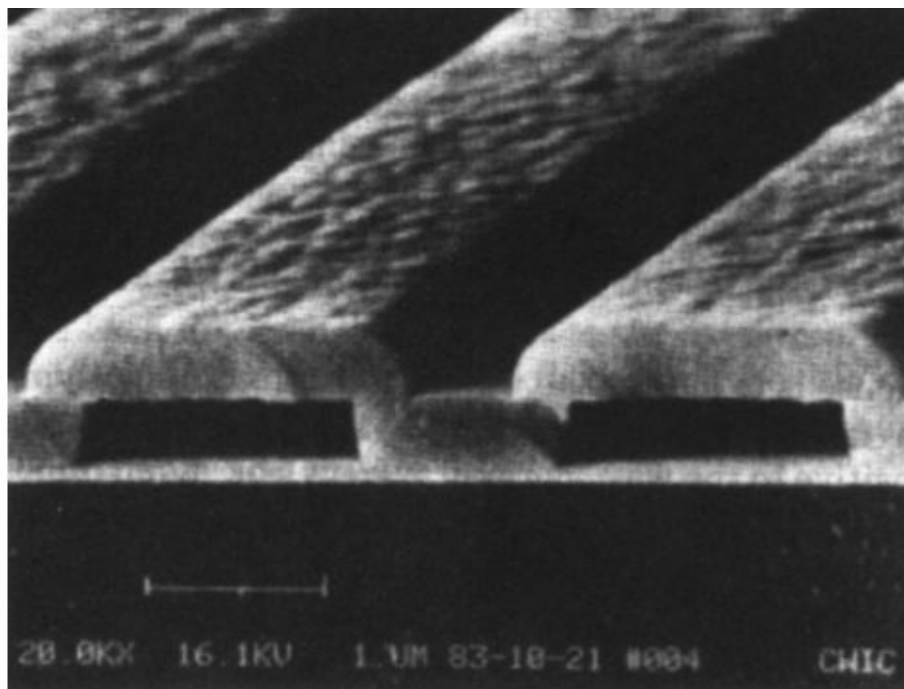


FIGURE 14.14 SEM declination of device cross-section.

Spreading Resistance Probe

Spreading resistance is a technique also used for measuring the junction depth. It is measured by a two probe on a beveled section of the wafer. As the probes are moved down to pass through the junction, they sense the change in conductivity type (N or P). This information, when plotted on the profile curve, also gives the junction depth ([Fig. 14.15](#)).

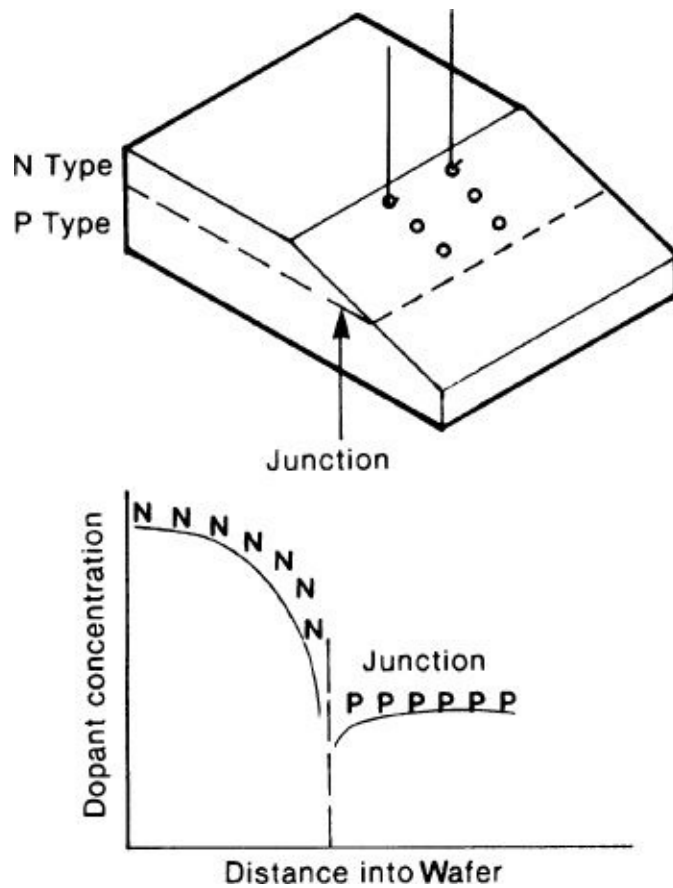


FIGURE 14.15 Spreading resistance.

Secondary Ion Mass Spectrometry

A secondary ion mass spectrometry (SIM) instrument bombards the wafer surface with a high-energy source that is a form of an ion-sputtering process. The collisions produce secondary ions from species on the surface that are a signature of the type of material. It can detect the dopant(s) in the junction region. The junction is located where the dopant type disappears from the stream of secondary ions.

Scanning Capacitance Microscopy

One of the options with atomic force microscopy (AFM) is a probe that measures capacitance. Since the capacitance of a doped layer changes as the junction is approached, it can be used for junction-depth detection. Sample preparation is the same as for spreading resistance measurements.

Scanning capacitance microscopy (SCM) offers the advantage of nanometer precision⁶ and may be the concentration profile and junction-depth measurement technique for submicron junctions.

Carrier Illumination Junction Depth

Carrier illumination is a nondestructive junction-depth measurement technique particularly useful for shallow junctions formed by ion implantation.² It starts with a laser “dot” (2 μm) over the junction area ([Fig. 14.16](#)). As the beam passes into the wafer and through the junction there is a buildup of excess carriers at the junction. A second laser beam is reflected off the excess carrier region. From there the reflected signal is analyzed and the junction depth determined.

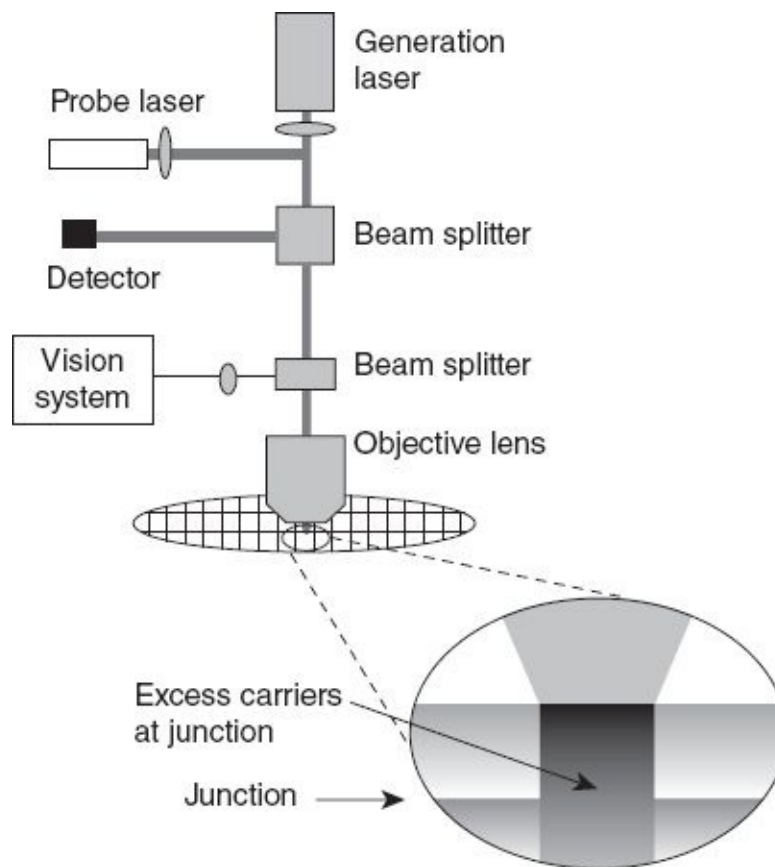


FIGURE 14.16 Noncontact Carrier Illumination™ measurement system.

This system can measure junction depths on monitor wafers or directly in a production wafer and can be installed as an in-line process control measurement.

Critical Dimensions and Line-Width Measurements

The exact dimensions required of each component in the circuit are controlled and influenced by *all* processes. Vertical dimensions are set by the doping and layering processes. The horizontal surface dimensions are produced in the lithography

process. As part of that process, the critical dimensions are measured at both the develop inspection and final inspection with microscopes.

Techniques include low-tech microscope image shearing, SEM, scatterometry, atomic force microscope (AFM), and device electrical measurements.

Optical Image-Shearing Dimension Measurement

An image-shearing attachment on a microscope is another method of critical-dimension measurement. A control on the unit allows the operator to optically separate (shear) the pattern into two images. To start the measurement, the two images are butted against each other ([Fig. 14.17](#)), then moved until the sheared images exactly overlap. The difference between the starting and ending values is the width of the pattern. Like the filar unit, an image-shearing unit must be calibrated to a standard. This method, while inexpensive and convenient to use is limited to widths greater than 1 μm .

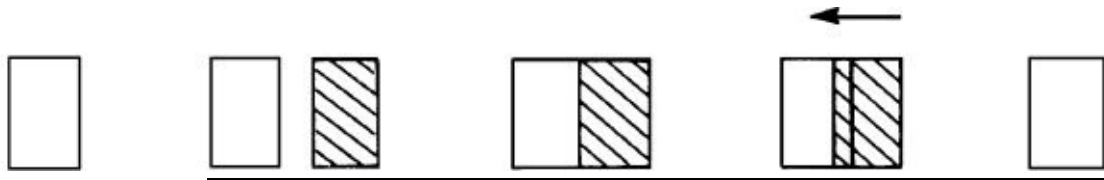


FIGURE 14.17 Single-image shearing.

Shape Metrology and Optical Critical Dimension

The SEM inspection tool described below is also used for very accurate measurements of line widths. In the nano-era, especially when using copper metallization, there is also interest in knowing and controlling the cross-sectional shape (3D shape metrology) of the hole or surface island ([Fig. 14.18](#)). This is accomplished with sophisticated SEMs that direct beams off the top surface, the sides, and the bottom surface to reconstruct the exact shape. Measurements are made directly on the wafer. A drawback is the optical critical dimension's (OCD)⁸ inability to measure isolated lines. The great goal is to have OCD systems integrated directly into the process tool, providing real-time measurement and process control ([Fig. 14.19](#)).

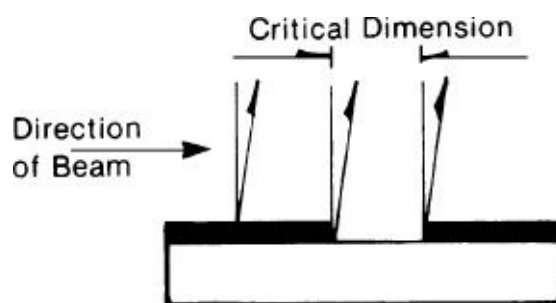


FIGURE 14.18 Reflectance CD measurement.

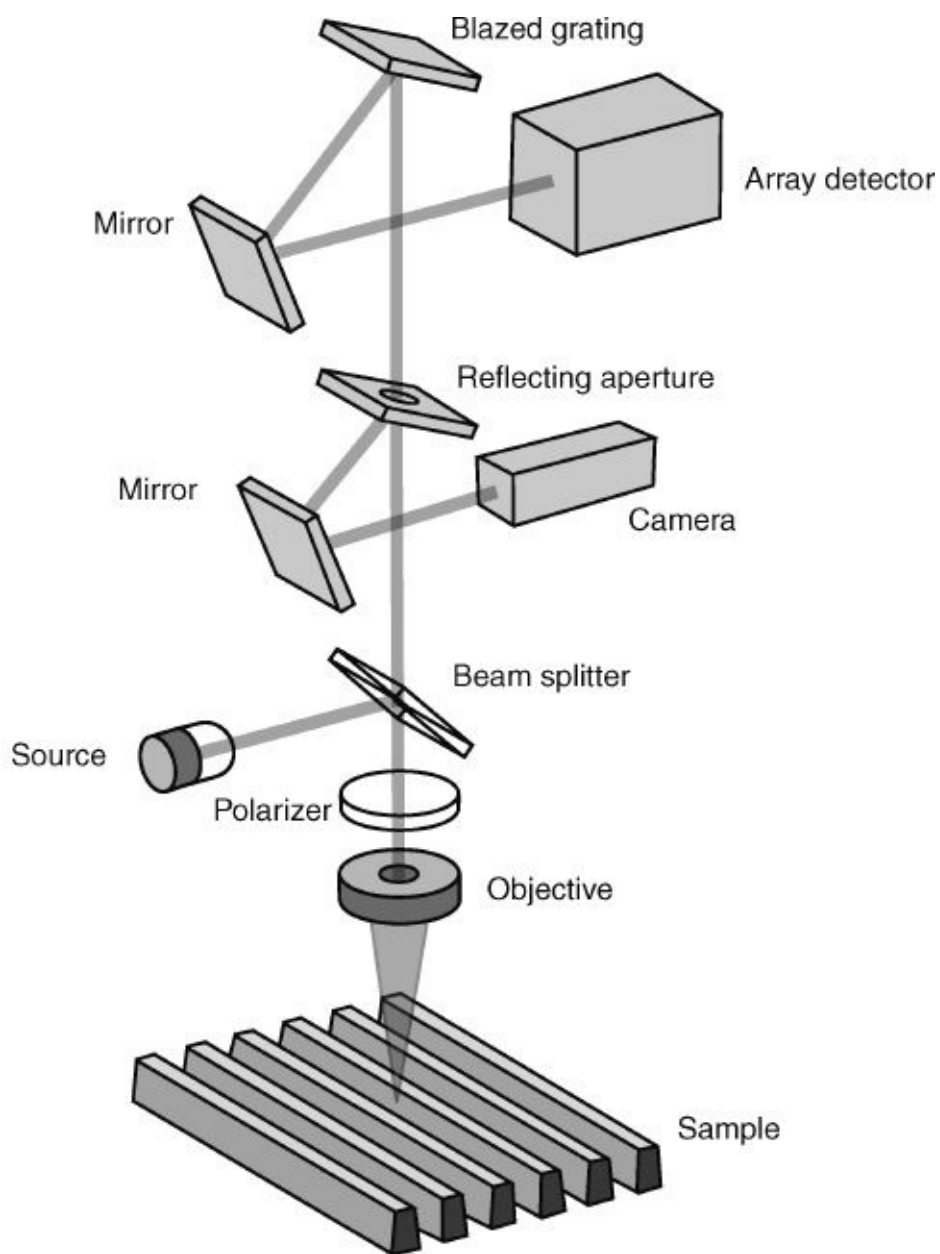


FIGURE 14.19 Schematic of OCD optics.

Contamination and Defect Detection

Detection of contamination and visual defects is essential to high yields and process control. Particulate contamination is detected primarily by visual techniques including high-intensity lights, microscopes, SEMs, and automatic machines. Chemical contamination is detected and identified by both Auger and electron spectroscopy for chemical analysis (ESCA) techniques. Like every microchip process the detection of killer defects and contamination has required enhanced techniques as device dimensions have gotten smaller.

1× Visual Surface Inspection Techniques

The first line of defense is a naked eye inspection of wafers, which in microscope terminology is 1× power. Operators quickly become sensitive to the way normal wafers look. Even minor changes in the surface appearance can identify a wafer or wafer batch for further inspection ([Fig. 14.20](#)).

Method	Visual contamination	Surface Defects	Alignment	Contamination Elements	Contamination Compounds	C.D.'s
1x Incident White Light	X	X				
1x Incident U.V. Light	X	X				
Microscope- Light field	X	X	X			
Dark field	X	X	X			
S.E.M.	X	X	X			X
Auger				X		
E.S.C.A.					X	
Filar						X
Image Shearing						X
Reflectance						X

FIGURE 14.20 Overview of surface-inspection techniques.

1× Collimated Light

The resolving power of the naked eye (1×) can be assisted by using a high-intensity white light, such as the beam of light from a slide projector ([Fig. 14.21](#)). Particulate

contamination is highlighted in the light beam when the wafer surface is viewed at an angle. The effect is similar to the highlighting of dust in the air by light streaming through a window.

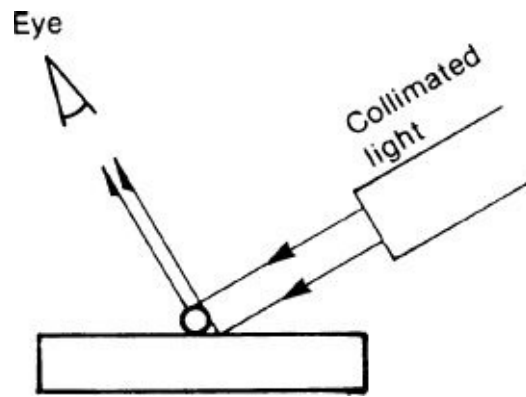


FIGURE 14.21 Collimated light inspection.

1× Ultraviolet

In actuality, the eye cannot see ultraviolet light, but ultraviolet light from a mercury-vapor lamp emits blue, green, and even some red light. Because ultraviolet is harmful to the retina, a filter is frequently placed over the light source to block out the ultraviolet. The primary benefit of the ultraviolet lights used in fabrication areas is that they are very bright, which means that the intensity of the scattered light is greater, therefore increasing the detection of surface contamination.

Microscope Techniques

Light-Field Microscope

The metallurgical microscope is the traditional workhorse of surface inspection. Shrinking dimensions and defect sizes have pushed standard microscopes to their resolution limit. But they are a fast and economical way to inspect for defects and contamination. And there are techniques to improve microscope resolution capabilities. The term *metallurgical* differentiates it from the standard microscopes found in biology labs. A biological microscope illuminates transparent samples by shining the light up through the sample. In a metallurgical microscope, the light is passed down to the nontransparent sample through the microscope objective ([Fig. 14.22](#)). The light reflects off the sample surface and is transmitted back up through the optics to the eyepieces. With white-light illumination, the picture in the field of view exhibits the surface colors, which helps identify particular components on the wafer.

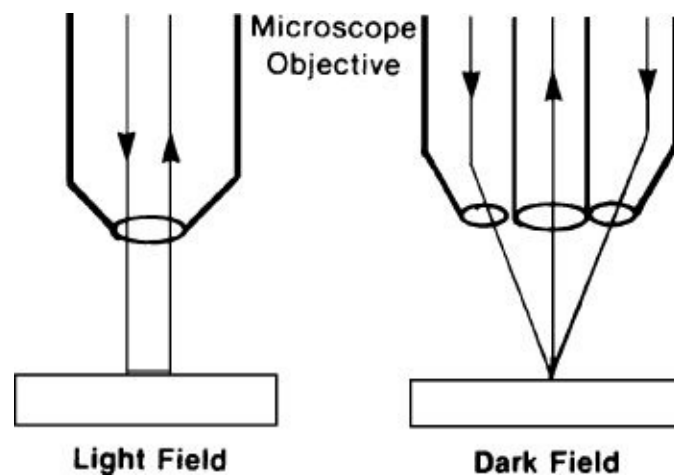


FIGURE 14.22 Light-and dark-field inspection.

A typical fabrication microscope is fitted with 10× or 15× eyepieces and a range of objectives from 10 to 100×. Increasing the total viewing power (eyepiece power times the objective power) reduces the field of view. This reduction requires more inspection time for the operator to look at the required sample inspection area on the wafer. The consequence is a slower inspection process. A trade-off power level, when inspecting LSI and VLSI devices, is 200 to 300× magnification.

The industry typically uses a microscope procedure requiring a sample inspection of several specific locations on the wafer. This procedure is easily automated with motorized stages. Most of the automated microscope inspection

stations feature automatic wafer placement on the stage and automatic focusing. Obviously, a microscope inspection procedure is used to judge surface and layer quality and (in masking) pattern alignment.

Just as image resolution in a patterning process is limited by the wavelength of light, so is bright-light inspection. With a broadband white-light source, the theoretical resolution limit is $0.30\text{ }\mu\text{m}$.⁹ Use of UV light sources and image processing can bring the resolution limit down to 80 nm.

Dark-Field Inspection

Dark-field illumination is achieved by fitting a metallurgical microscope with a special objective ([Fig. 14.22](#)). In this objective, the light is directed to the wafer surface through the outside of the objective body. It impinges on the surface at an angle and passes up through the center of the objective. The effect on the “picture” in the eyepieces is to render all flat surfaces black. Any surface irregularities, such as a step or pieces of contamination, appear as bright lines. Dark-field illumination is more sensitive than light-field to any surface irregularity. It has the drawback of limiting the ability to discern the nature of the surface irregularity. A passable surface dimple may look the same as a rejectable piece of contamination.

Dark-field resolution of defects is enhanced with the use of laser light¹⁰ sources and multiple sources.

Confocal Microscopes

Distinguishing fine details with a normal microscope is hampered by stray light rays interfering as they bounce off the various surfaces and from the different plane depths on a wafer surface. Confocal light sources minimize the scatter and provide greater detail by limiting the returning light rays to a narrower plane. This is done by passing an intense white light or laser beam through holes in a disk (positioned between the light source and the wafer), spinning at high rpm. Only light rays from the inspection plane of interest bounce back through the holes. Confocal systems are capable of imaging submicron dimensions.⁴ Laser beams used in a confocal system cause the sample to fluoresce (glow). The resultant image, after x-y scanning and computer processing, has resolution in the nanometer range and can display layers of registered translucent samples.

Other Microscope Techniques

Optical technology is capable of providing many evaluation techniques beyond simple light-and dark-field viewing, such as phase contrast and fluorescence microscopes. Each allows the viewer to determine additional visual information about the surface. Phase contrast brings out surface irregularities in the vertical plane, and fluorescence-illuminated microscopes use ultraviolet illumination sources. In the ultraviolet light, organic residues (photoresist, cleaning chemicals) not easily visible in white light are brought into view. Their use and interpretation generally require technicians trained beyond the level of production operators.

Scanning Electron Microscopes

Conventional optical microscopes are limited in their ability to provide accurate information about the wafer surface. First, their resolving power is limited by their optical light source. The ability of a viewing system to distinguish detail is related to the wavelength of the light (radiation used). The shorter the wavelength, the smaller the detail that can be seen.

Depth of field is another viewing factor. It relates to the ability of the system to keep two planes in focus simultaneously. A conventional photograph, with the subject in focus and the background out of focus, has a background beyond the depth of field limit of the camera. In a microscope, the depth of field decreases as the power (magnification) of the system is increased. If the power is increased to see the surface “closer,” the operator may not be able to see the top and bottom surfaces in focus. Constant refocusing results in loss of information and a longer inspection time.

Magnification is the third limiting factor of optical microscopes. An optical system with white-light illumination is limited to about 1000× magnification with conventional objectives. The oil-immersion technique pushes the limit up, but it is unacceptable, because it is too slow, too messy, and a possible source of contamination to the wafer.

All three limitations are overcome by using a scanning electron microscope. The microscope varies from an optical one in many aspects. The “illumination” source is an electron beam scanned over the wafer or device surface. The impinging electrons cause electrons on the surface to be ejected. These secondary electrons are collected and translated into a picture of the surface ([Fig. 14.23](#)) on either a screen or a photograph.

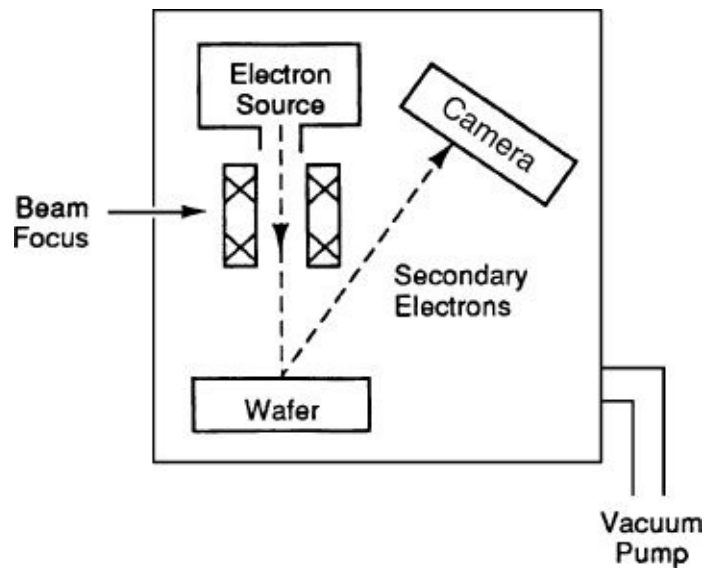


FIGURE 14.23 SEM analysis.

Scanning electron microscope (SEM) analysis requires that the wafer and beam be in a vacuum. The electron beam has a much smaller wavelength than white light and allows the resolution of surface detail down to submicrometer levels. It eliminates depth of field problems; every plane on the surface is in focus.

Magnification is similarly very high, with a practical upper limit of 50,000 $\times$. A tilting wafer holder in an SEM allows the viewing of the surface at angles, which enhances the three-dimension perspective ([Fig. 14.24](#)). Surface details and features can be viewed at advantageous angles.

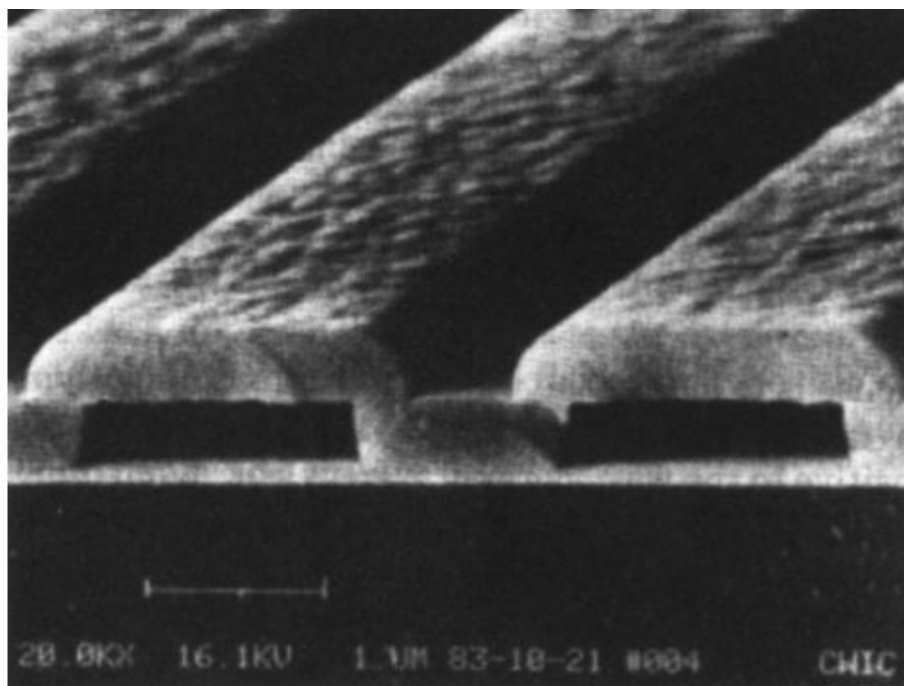


FIGURE 14.24 SEM declination of device cross-section.

Some materials, such as photoresist, do not give off secondary electrons in response to e-beam bombardment. To inspect a photoresist layer in an SEM, the resist layer is covered with a thin layer of evaporated gold. The gold layer conforms to the topography of the photoresist layer. Under e-beam bombardment, the gold gives off secondary electrons, thus resulting in an SEM picture that is an accurate reproduction of the underlying photoresist layer.

Transmission Electron Microscope

A standard SEM instrument has a resolution range of 20 to 30 Å,¹² which presents a barrier for inspection of ULSI devices. However, passing (transmission) the electron beam through a thin sample increases the resolution to 2 Å. As attractive as the increased resolution may be, sample preparation is time-consuming and requires precision. The transmission electron microscope (TEM) is best used as an off-line laboratory tool.

This principle is used to measure film thicknesses. In an SEM, the energetic e-beam causes X-rays to leave the surface. Fortunately, the X-ray energies are reflective of the materials (elements) presence on the surface and can be analyzed for additional information. This is accomplished when an X-ray spectrometer is added to the SEM to detect dispersed X-rays. This additional information is about the chemistry of the materials.

Optical Profilometry

This is a microscope type instrument that divides the light beam in two with a beam splitter. One beam is reflected off a reference mirror. Combining the reflected beams causes interference that in turn provides surface topography and contours. It also can determine step heights and provide 3D images.¹³

Capacitance-Voltage Plotting

Mobile ionic contamination (MIC) is electrically active ions that have been incorporated into the device during processing. They come from contaminated process materials, dirty process stations and chambers, and environmental contaminants. They are not detected with surface inspection methods. However, they do show up in the electrical measurements of transistors and diodes. The various parts of a semiconductor device are made from precision doping. MICs act in the device as a dopant and change the electrical characteristics of the device. A common electrical test is capacitance-voltage plotting on MOS transistors. The technique is described in the section on electrical device testing.

Automated In-Line Defect Inspection Systems

Automatic Defect Detection

The detection of ever smaller-sized particles on the wafer surface has led to the use of laser beam and e-beams as the detection illumination. Two advantages accrue from the use of these beams. First is the obvious ability to detect smaller particle sizes ([Fig. 14.25](#)) due to the high intensity of the reflected light, which makes very small surface particles detectable (helium and neon lasers are the usual sources). E-beams have even smaller spot size and are the system of choice for nano-sized patterns, especially for copper or low-k dielectric metallization systems.

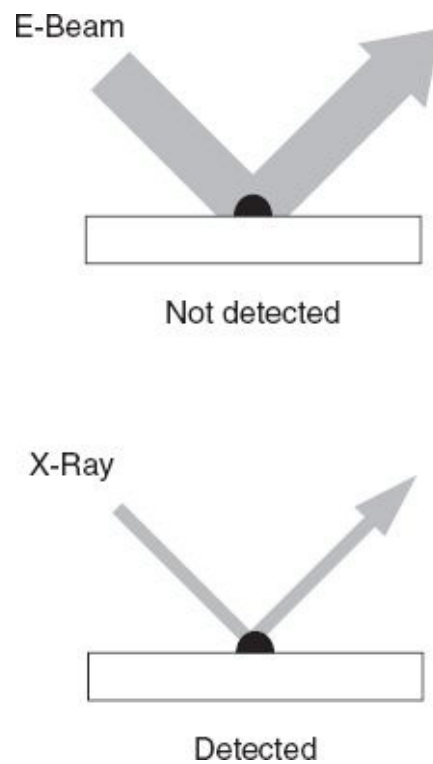


FIGURE 14.25 X-ray detection of small surface particles.

The second advantage is automation. Laser inspection equipment is easy to automate such that the inspection is automated from cassette to cassette, and the quantity and size of the contamination on the surface can be determined. The information is used to produce a map of the wafer surface showing the size, location, and density of surface contamination. Plus the inspections occur in-line in real time and the onboard computer analysis can detect patterns and deviations from specifications. Thus production and process corrections can be made quicker and with better certainty. These last features are the basis for “factory” process control.

Some systems mix bright-and dark-field viewing and process the information

into maps or databases ([Fig. 14.26](#)). Advances in image processing have allowed automatic defect-and pattern-distortion-detection instruments. The machines mate image processing and computer technology. They have scanners with laser or optical light sources that move over the wafer surface. In one version, the computer is preprogrammed with the design pattern for the circuit. In the *die-to-database* system, each die is scanned and the results compared to a stored image of the mask or reticle for the particular layer. The scanner looks for added or missing parts from the expected pattern. Anything missing from or on the die that is not in the database is flagged for further inspection. The presumption is that, if an image is not in the database, it is probably a defect of some kind. Defect locations are recorded and surface maps printed out. Engineers can go back to the wafer or mask and determine the exact nature of the problem.

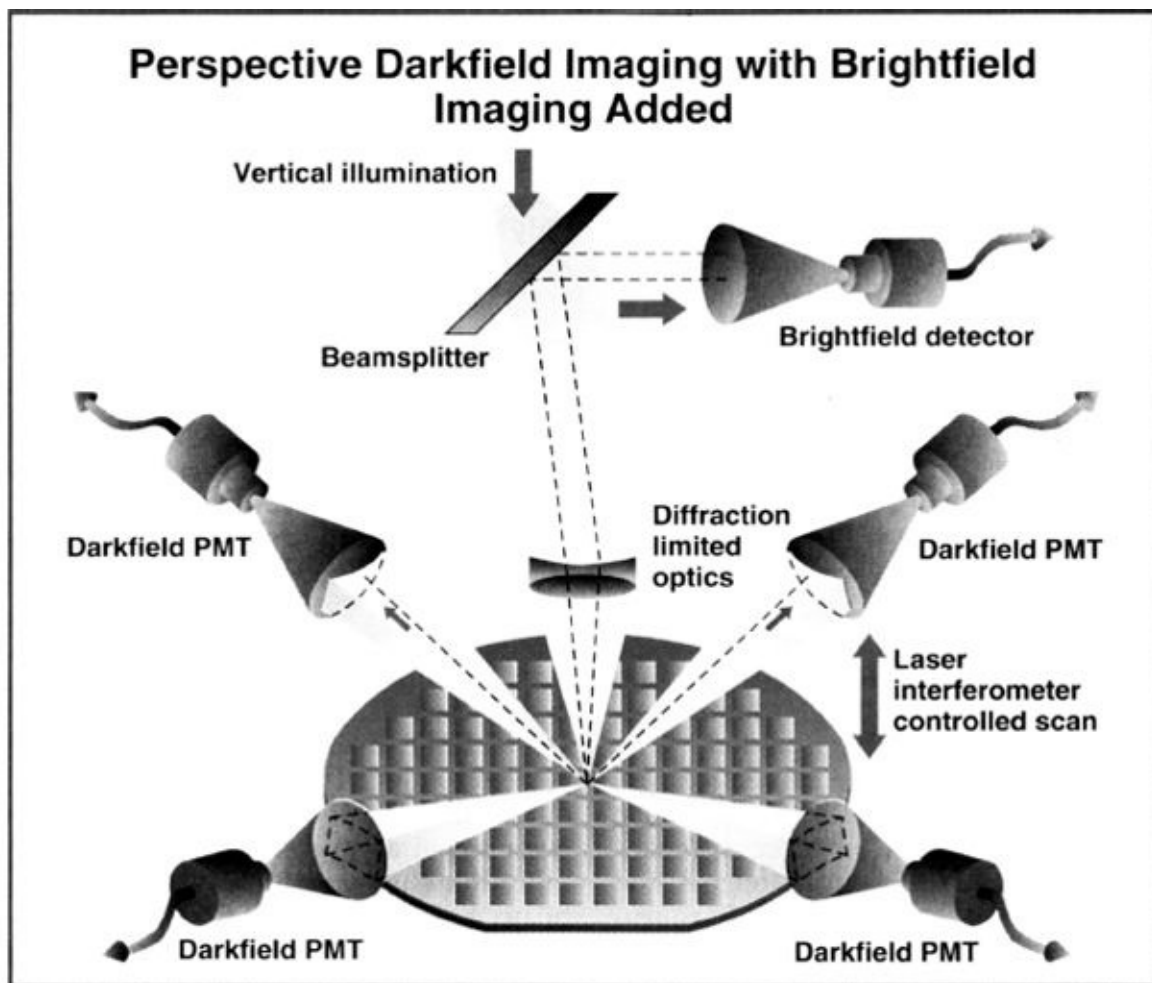


FIGURE 14.26 Mixed dark-field and light-field system.

Another system, called die-to-die inspection, compares the adjacent die on the wafer or mask. A die is scanned and the pattern recorded in the computer memory. A second die is also scanned, and any deviations between the two are recorded. This

system will not detect any repeated pattern defect that occurs on every die, but it will pick up random defects that have a very low probability of occurring in exactly the same spot on two adjacent die. In both machine types, the information from the surface is captured by charge-coupled device (CCD) cameras or photomultiplier tubes.¹⁴ For inline inspection, a grave concern is calibration and standardization.

In addition to onboard verification of the electronics, most systems use a standard wafer to verify machine operation. Several inspection processes are used. While an automatic machine can find defects, the determination of which ones are “killer” defects is more difficult. With overall defect counting, it is possible that the machines may report a rise in defect count, yet the increase may contain only minor, non-killer types. Verification by humans is still a critical part of any defect-detection and management system.

Another advance in automated wafer-surface inspection is combining different techniques in the same inspection tool. One example houses scatterometry, ellipsometry, reflectometry, along with topological analysis to provide a comprehensive evaluation of the surface condition. Defects, thickness uniformity, and surface conditions can be determined in one inspection station.¹⁵

The development of copper or low-k dielectric combinations has introduced a whole new host of defects to the process.¹⁶ [Figure 14.27](#) lists a sampling of defects associated with these metal systems.

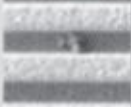
Black spots		Post-Cu CMP	Slurry/low-k Rxn
Etch residue		Post-trench etch	Post-etch clean
Low-k bumps		Post-low-k dep	Cu hillocks
Craters		Post-via etch	Etching of "bumps"
Cracking		Post-via etch	Stress of ILD
Resist poisoning		Via photo	Via patterning
Bond pad blistering		Low- dep	Adhesion loss

FIGURE 14.27 Sources of low-k defects. (Courtesy of Semiconductor International.)

General Surface Characterization

Atomic Force Microscopy

Several of the measurement techniques are multiuse, such as SEMs. A newer entry in the general category is *atomic force microscopy* (AFM). It is a surface profiler that scans a delicately counterbalanced probe over the surface ([Fig. 14.28](#)). Probe and surface separation is so small (about 2 Å) that atomic forces between the surface and probe materials actually affect the probe. Sophisticated electronics measure the relative position of the probe as it moves across the surface, thus compiling a three-dimensional map of the surface¹⁷ ([Fig. 14.29](#)).

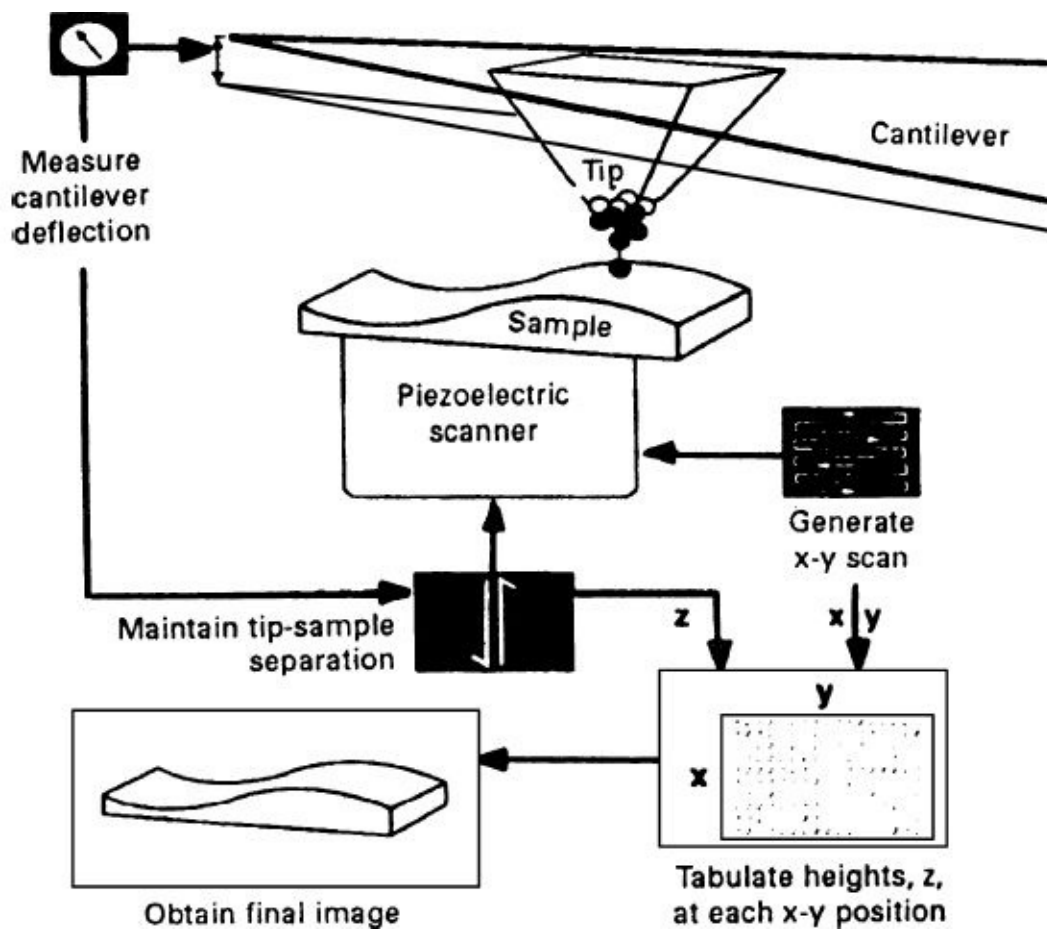


FIGURE 14.28 Atomic force microscope (AFM). (Source: Semiconductor International.)

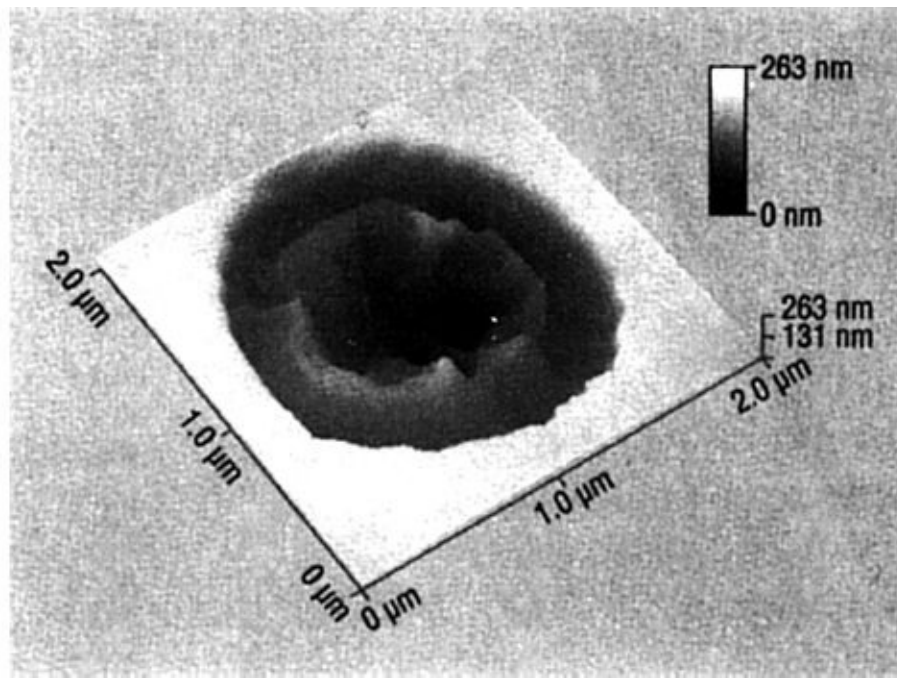


FIGURE 14.29 Atomic force microscope (AFM) surface map.

AFM sensitivity is in the 1-Å range and can be operated in contact or noncontact modes. The first AFM production instrument was introduced in 1990. Projections are for it to become an integral part of the inspection arsenal. By its nature, AFM can characterize grain sizes, detect particles, measure surface roughness, and provide critical dimension measurements—all in three dimensions. In one novel use, an AFM probe has been combined with an optical microscope objective.

Scattrometry

The search for fast, accurate, and nondestructive surface inspection tools led to the *scatterometry* metrology. Accuracy in optical systems is limited by the wavelength of the light. Roughly, particles or surface features smaller than the wavelength used cannot be detected. However, a scattered beam can give information about surface features smaller than the wavelength. A scatterometry system has the wafer at the center of curvature of a screen (Fig. 14.30). An incident laser beam is scanned over the surface and is reflected and “scattered” from the surface onto the screen. A camera with a microprocessor captures the screen image to reconstruct the surface that produced the particular pattern on the screen. Like AFM, this technique has the potential of measuring grain sizes, contours, and CDs. It can also measure latent images in undeveloped photoresist and characterize phase-shift masks.

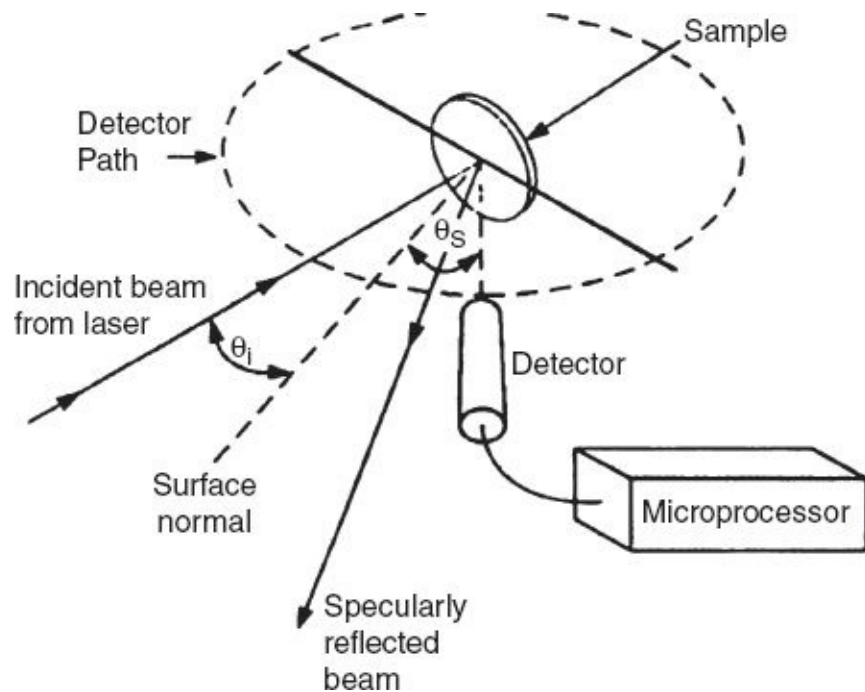


FIGURE 14.30 Arrangement of a scatterometer. (Source: *Solid State Technology*, March 1993.)

Contamination Identification

Pushing into the nano-era is requiring more detailed information about the nature of the various contaminants that end up on the wafer surface or in the deposited layers. Contamination types, forms, amounts, and other data are needed to maintain clean processes and products. In this section, the instruments used to collect this data are examined. In general, all of the techniques are based on a common phenomenon. Whenever a surface is excited with energy, there will be energy given off. The energy given off will have characteristics indication the material(s) on the wafer surface.

Auger Electron Spectroscopy

In an SEM, a range (spectrum) of secondary electrons is released by the impinging electron beam. One portion of this spectrum is the electrons that are released from the top several nanometers of the surface. These electrons, known as *Auger electrons*, have energies characteristic of the element that emits them. Thus, sodium and chlorine each give off different Auger electrons.

The collection and interpretation of Auger electrons allows the identification of the surface materials, including contamination. In operation, the e-beam is scanned across the wafer. The ejected Auger electrons are analyzed for their energies (wavelengths) and printed out on an x-y plotter ([Fig. 14.31](#)). Energy peaks at specific wavelengths indicate the presence of specific elements on the surface.

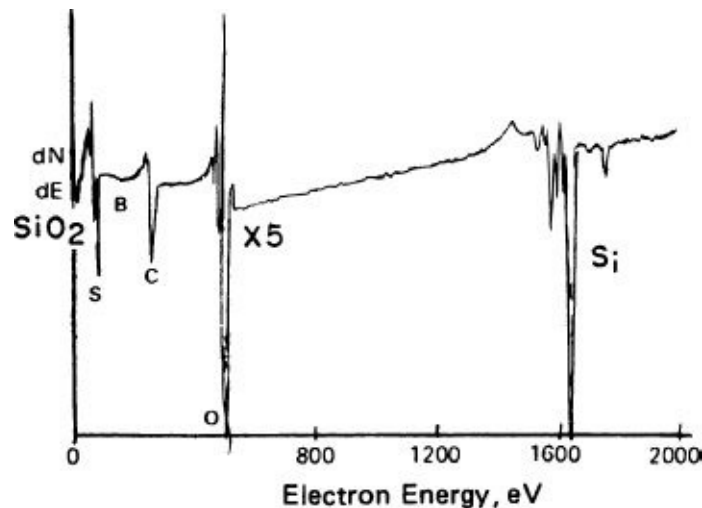


FIGURE 14.31 Typical Auger trace.

Scanning Auger microanalysis (SAM) is limited to the identification of elements. This technique cannot identify the chemical state of the element, what it may be combined with, or its quantity on the surface. Salt (NaCl) contamination is identified only as the presence of sodium and chlorine on the surface.

Electron Spectroscope for Chemical Analysis

Solving surface-contamination problems often requires knowledge of the chemical state of the contaminant. An Auger detection of chlorine on the surface does not reveal whether the chlorine is present as hydrochloric acid or a trichlorobenzene. Knowledge of the form of the chlorine expedites locating and eliminating the process source of contamination.

The electron spectroscope for chemical analysis (ESCA) is an instrument used to determine surface chemistry. The instrument works on principles similar to the Auger technique. However, X-rays instead of electrons are used for the bombarding radiation. Under bombardment, the surface gives off photoelectrons. Analysis of the photoelectron information leads to the determination of the chemical formula of the contamination. Unfortunately, the ESCA X-ray beam is wider than most integrated circuit features. The beam diameter limits the technique to a macro-surface analysis. By contrast, an Auger electron beam can zero in on specific bits of contamination.

Time of Flight Secondary Ion Mass Spectrometry

TOF-SIMS surface analysis can use a number of incident radiation sources, such as an Nd YAG laser, a Cs ion beam, or a Ga ion beam.¹⁸

For this method, the impinging energy pulses and ionizes surface material and produces secondary ions that are accelerated to a mass spectrometer. In the spectrometer, each ion's time of flight from the surface is measured. The time of flight is indicative of the species on the surface. TOF-SIMS samples material depths of only several tenths of a nanometer.¹⁹ Along with the ability to characterize both organic and inorganic contaminants, it is a mainstay of an analytical lab.

Evaluation of Stack Thickness and Composition

The top of the wafer is becoming more and more complicated with stacks of various materials. Each material can be measured for thickness and composition from monitor wafers included in the deposition process. However, their effect in and on the circuit comes from the combination of thicknesses and compositions. A combination machine to determine these parameters, *in situ*, uses an electron beam and X-ray spectrometers ([Fig. 14.32](#)). An e-beam is directed at the stack, which in turn kicks off X-rays. A series of X-ray spectrometers arrayed around the sample analyze the rays. Because each thickness and material composition creates a unique X-ray signature, the machine's onboard computer can run an analysis and determine the thickness and composition of each component in the stack.²⁰

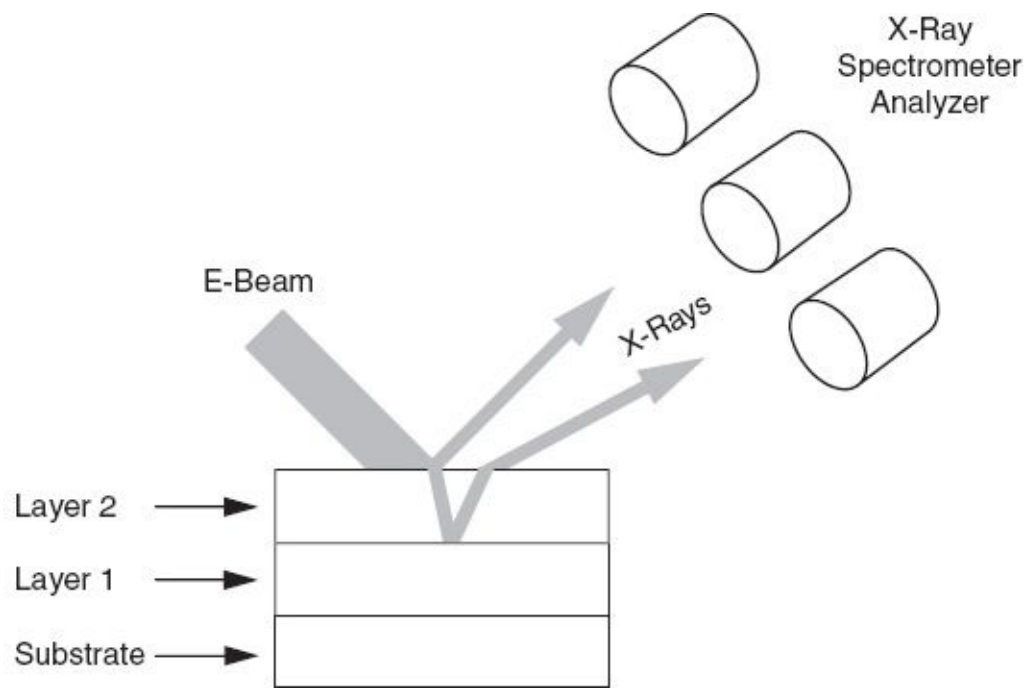


FIGURE 14.32 E-beam or X-ray detection of material stacks.

Device Electrical Measurements

During the process, it is necessary to make direct measurements of the actual device parameters. These measurements are usually made on special devices in the test die or on special structures in the scribe line areas.

A great deal of information about the process is available from these electrical tests. In this section, we explain the basic tests and several of the more obvious and frequent device failures. Often, these tests are referred to as *parametric testing*. The reference is to the measurement of the device electrical parameters as opposed to testing the overall function of the device or circuit, which occurs at wafer sort at the end of the fab process. A complete device and process troubleshooting guide is beyond the scope of this text.

Equipment

The basic equipment required to perform device electrical tests are (1) a probe machine ([Fig. 14.33](#)) with the capability of positioning needlelike probes on the devices, (2) a switch box to apply the correct voltage, current, and polarities to the device, and (3) a method of displaying the results. If many devices are to be measured, the probes may be fixed on a printed circuit board-like structure called a *probe card*. They are the same cards used in the wafer-sort process. Display may be by a curve tracer or by a special digital voltmeter. A curve tracer is a specially designed oscilloscope set up to display voltage on the horizontal scale and current on the vertical scale. They show the characteristic voltage or current traces for each type of measurement. Example displays are shown in the following sections for each test. In high-production situations, the tests are performed digitally and the results displayed on a read-out. These measurement systems have the capability of storing and correlating data.

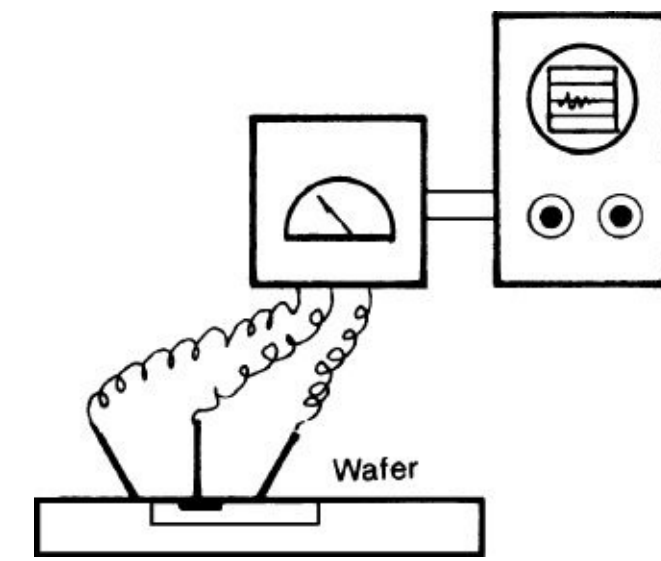


FIGURE 14.33 Device-measurement equipment.

In advanced systems, the probe station may be automated to sequentially test several die. The switching can also be automated to allow the equipment to perform a series of tests in a predetermined sequence. Automated systems can display graphs of the test results as well as analysis of the data. The individual tests will be explained, as performed on manual equipment. The shape and relationship of the traces displayed on the oscilloscope are very helpful for understanding device operation.

All the device measurements are made in basically the same way. A voltage is

applied to the component contact probe, and the resultant current flowing between the contacts is measured. The results are displayed on the oscilloscope screen or display screen. The shape of the trace is governed by the dimensions of the device, the doping, and the presence of junctions. The current is modified by the resistance of the component or the presence of junctions. Properly working components will display known traces. The fundamental relationship is that of Ohm's law: $R = V/I$

Resistors

Resistor measurements are made by contacting each end of the resistor and applying a voltage. Current passage between the probes will be governed by the nature of the material or, in electrical terms, its resistivity. During the test, the applied voltage is varied from zero to some higher value. On the screen, the voltage value is displayed on the x axis, with current measured on the y axis ([Fig. 14.34](#)).

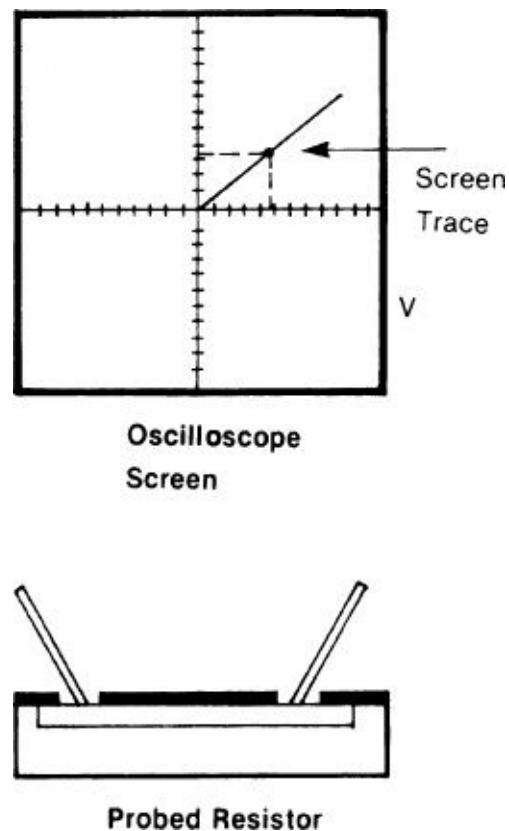


FIGURE 14.34 Resistor.

The value of the resistance is calculated by dividing one of the voltage values by the corresponding current value. One might ask, “Why not determine the resistance by simply measuring the voltage and current values with a meter?” In other words, why display the values on the screen? The answer lies in the quality of information gained from the trace. A resistor’s V/I relationship should be a linear one (straight line). Any deviation from linearity could indicate a process problem, such as high contact resistance or a leaking junction.

Diodes

Diodes function as switches in a circuit. This means that a diode can pass current in one direction (forward bias) and not in the other (reverse bias). Checking diode operation in either direction requires measuring the diode with the proper polarities, as shown.

As the voltage is increased in the forward direction, current immediately starts flowing across the junction and out of the diode ([Fig. 14.35](#)). The initial resistance to that flow comes from contact resistance and a small resistance at the junction. After the resistances are overcome, there occurs a “full” flow of current through the diode. A diode is designed to have this condition occur at some minimum voltage value. This voltage is called the *forward voltage*. If the diode forward voltage value exceeds the design value, it is out of specification.

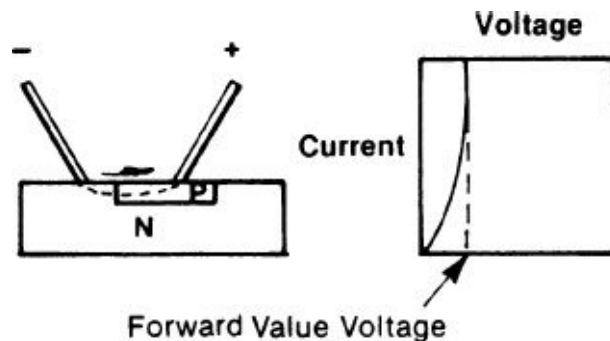


FIGURE 14.35 Diode forward bias.

In the reverse direction, the diode is designed to block current flow as long as the voltage stays below a specified value. In the reverse condition, a small current, called a *leakage current*, always flows across the junction ([Fig. 14.36](#)). Increasing the voltage increases the leakage current to a level at which the junction breaks down, allowing “full” current flow. This value of the voltage where full flow begins is called the *breakdown voltage* ([Fig. 14.37](#)). Circuits are designed to operate at a voltage level below the designed breakdown voltage of the diode to take advantage of the blocking nature of the junction. Lowered breakdown voltages are generally the result of processing mistakes or excessive contamination.

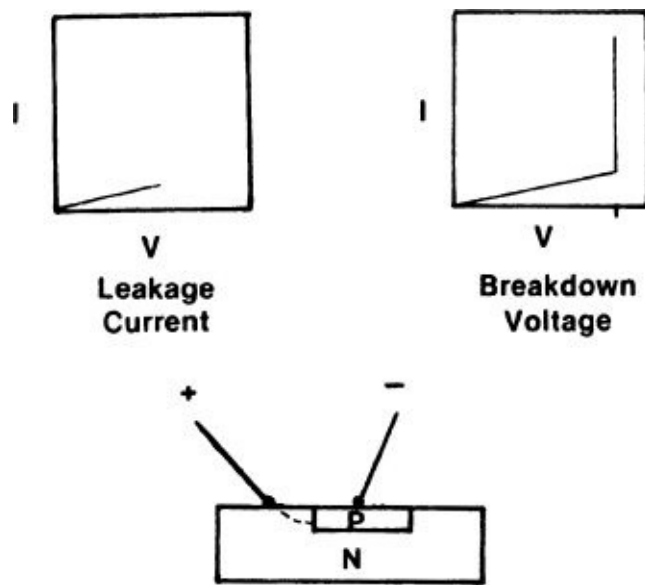


FIGURE 14.36 Diode reverse bias measurement.

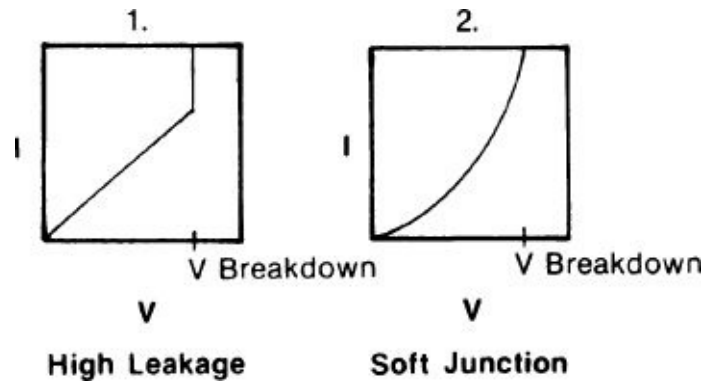


FIGURE 14.37 Junction leakage.

Exceeding the breakdown voltage does not (normally) permanently damage the junction. However, if the applied voltage is extremely high, the diode (junction) can sustain permanent physical damage from a high current flow.

A second parameter determined during this test is the current at the breakdown voltage. A small amount normally occurs as illustrated below. Contamination and/or improper processing can result in additional leakage current.

Trace 1 in [Fig. 14.37](#) shows a diode with a small amount of leakage. The amount of current increases as the voltage is increased. Eventually, the breakdown voltage is reached, and the diode becomes fully conducting. In trace 2, gross leakage is demonstrated. The junction leaks current with every increase in voltage, and the problem is so severe that a breakdown level is never reached, and the diode never operates as a current block.

Bipolar Transistors

Bipolar transistors, as explained in [Chap. 16](#), are three-region, two-junction devices. Electrically, they can be thought of as two diodes back to back. Many tests are performed to characterize bipolar transistors. The individual junctions are characterized separately and the whole transistor operation measured. The junctions are probed for forward and reverse characteristics. The breakdown voltage (BV) tests are followed by the probing of the entire transistor.

Individual junction probe measurements are designated by the letters BV followed by lowercase letters that indicate the particular junction being probed. BV_{cbo} indicates the breakdown voltage measured between the collector and base regions. The “o” indicates that the emitter region is “open”—it has no voltage applied to it. BV_{ceo} is the breakdown measured between the collector and emitter.

The forward voltages of the two bipolar structure junctions are also measured. V_{be} is the forward voltage of the emitter-base junction, and V_{bc} is the forward voltage of the base-collector junction. BV_{iso} probes the collector-isolation junction for leakage currents.

A principal electrical measurement of a bipolar transistor is the beta ([Fig. 14.38](#)) measurement. This is a measurement of the amplification characteristic of the transistor. In a bipolar transistor, the current flows from the emitter to the collector, through the base (see [Chap. 16](#)). The base current is varied to change the resistance in the base region. The amount of current flowing out of the collector (from the emitter to base) is regulated by the base resistance.

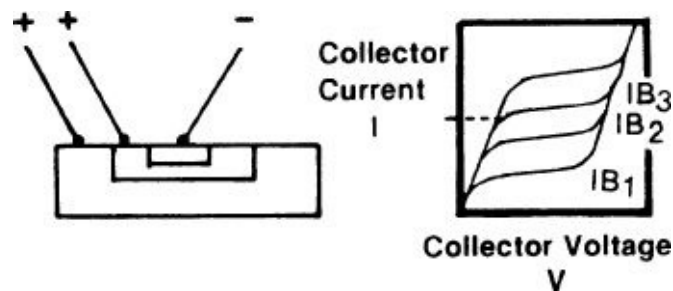


FIGURE 14.38 NPN transistor beta measurement.

The amplification of the transistor is defined as the collector current divided by the base current. This number is known as *beta*. Thus, a beta of 10 means that a 1-mA (milliampere) base current will give rise to a 10-mA collector current. The beta of a transistor is determined by junction depths, junction separations (base width), doping levels, concentration profiles, and a host of other process and design

factors. Measurement of beta is done by a variation of the BV_{ce0} measurement. A BV_{ce0} measurement at a specific base current is performed. In this mode, the emitter base junction is forward biased.

The collector characteristic of the transistor is displayed on the oscilloscope screen. The almost horizontal lines represent increasing base current values (IB_1 , IB_2 , and so on). With every increase in base current, a corresponding increase in collector current occurs. Calculation of the beta value takes place from the data displayed on the screen.

Collector current is determined from the vertical axis (dotted line). The base current is calculated by multiplying the number of horizontal lines (steps) by the scale value for each step (from the oscilloscope).

MOS Transistors

MOS circuits are also made up of resistors, diodes, capacitors, and transistors. The first three are measured by the same methods used to measure bipolar circuit components. Like the bipolar transistor, the MOS transistor is composed of three regions, in this case called the source, gate, and drain (see [Fig. 14.39](#) and [Chap. 16](#)). Measurement of this type of transistor consists of determining the reverse and forward values of the source and drain junctions. The functioning of the gate is determined by the threshold voltage test.

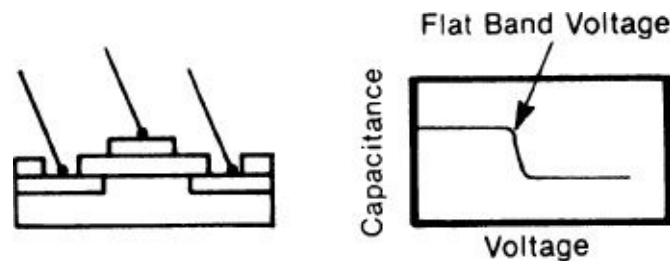


FIGURE 14.39 Threshold voltage measurement.

A MOS transistor has the source region forward biased. Because of the high resistivity of the gate region, the forward current does not reach the drain. A voltage applied to the gate at a specified level (threshold) will cause enough charges to appear in space between the source and drain and under the gate region to form a conducting channel that allows the source current to reach the drain region. Every MOS transistor is designed to operate at a specific threshold voltage. This value is measured using the capacitance-voltage technique. The gate voltage is continuously increased, while the capacitance of the gate structure is monitored.

A capacitor is a storage device. Initially, during the voltage-increase portion of the measurement, the capacitance does not change. At the threshold voltage level, the inversion layer forms and acts like a capacitor. Because two in-series capacitors have a combined lower capacitance than the sum of the two, the result is a drop to a combined lower capacitance. MOS transistors also exhibit amplification characteristics. The gain is defined as the source-drain current divided by the gate current. The source-drain characteristic for various gate currents is shown ([Fig. 14.40](#)).

$$\text{Gain} = \frac{\Delta I_{DS}}{\Delta V_{GS}}$$

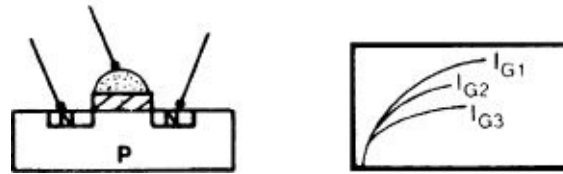


FIGURE 14.40 Gain characteristic of an MOS transistor.

Capacitance-Voltage Profiling

A variation of the threshold voltage test is used to test for the presence of mobile ionic contamination in the oxide. The test is performed on specially prepared test wafers. A thin oxide is grown on a “clean” silicon wafer. After oxide growth, aluminum dots are formed on the wafer by evaporation through a mask ([Fig. 14.41](#)). Dot evaporation is usually followed by an alloy step to ensure good electrical contact between the aluminum and oxide. Special MOS capacitors may be formed on the wafer in circuit test sites.

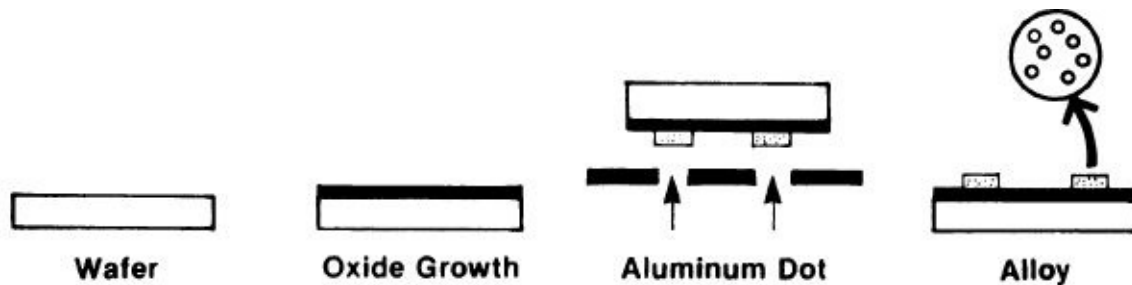


FIGURE 14.41 Preparation of C/V test wafer.

The “dotted” wafer is placed on a chuck, and a probe is placed on the aluminum dot. This structure is actually an MOS capacitor. A voltage is applied to the dot and gradually increased as the capacitance of the structure is simultaneously measured. The results are printed out on an x-y plotter with capacitance on the y axis and voltage on the x axis ([Fig. 14.42](#)).

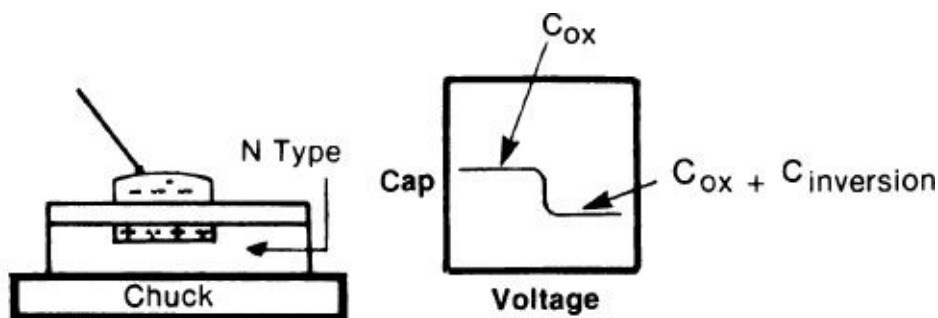


FIGURE 14.42 C/V plotting—first test.

At a voltage level known as the *threshold voltage* (or *inversion voltage*), charge starts to build up at the silicon surface. The charges “invert” the conductivity type from N-type to P-type. The inverted layer has a capacitance of its own. Electrically, the structure now has two capacitors in series. The total capacitance value of the two is less than the sum of the two by the relationship:

$$1/C_{\text{total}} = 1/C_{\text{oxide layer}} + 1/C_{\text{inversion layer}}$$

The trace on the x-y plotter drops vertically to the new capacitance level.

The second step in the process is to force the mobile positive ions in the oxide to the SiO_2 -silicon interface. This is done by simultaneously heating the wafer to the 200 to 300°C level and placing a positive 50-V bias on the structure (Fig. 14.43). The elevated temperature increases the mobility of the ions, and the positive bias “repels” them to the oxide-silicon interface.

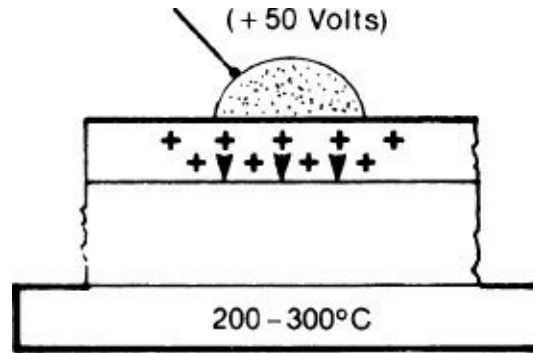


FIGURE 14.43 C/V profiling—ionic charge collection.

The last step in the process is a repetition of the initial C/V plot. However, as the voltage increases, inversion does not start at the same level as in the initial test (Fig. 14.44). The positive charges at the interface require additional negative voltage to “neutralize” them before inversion can happen. The result is a C/V plot identical to the original but displaced to the right. The additional voltage required to complete the plot is known as the *drift* or *shift*.

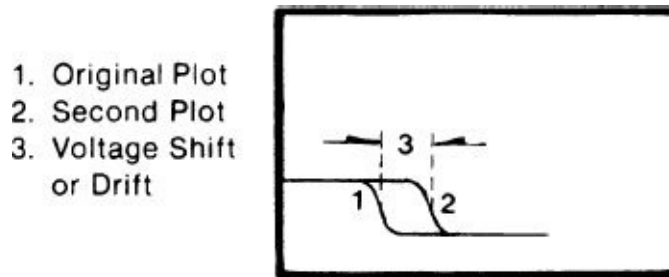


FIGURE 14.44 C/V re-profile.

The amount of the shift is proportional to the amount of mobile ionic contamination in the oxide, the oxide thickness, and the wafer doping. C/V analysis cannot distinguish the element (Na, K, Fe, and so forth) that was in the oxide—only the amount. Neither can this test determine where the contamination came from. It may have come from the wafer surface, any cleaning step, the oxidation tube, the evaporation process, the alloy tube, or any other process(es) the wafer has been through.

C/V analysis is usually a part of the evaluation of any process changes that may contaminate a wafer, such as a new cleaning process. To make the evaluation, C/V wafers are divided into two groups. One group receives normal processing as detailed above. The second group goes through the proposed cleaning process, usually between the oxidation and aluminum steps of the test. The drift on the standard processed wafers is compared with the drift of the experimental group. An increased drift on the experimental wafers would indicate that the proposed cleaning process actually contaminated the wafers with mobile ionic contamination.

Acceptable C/V drifts vary between 0.1 and 0.5 V, depending on the sensitivity of the device being made. C/V analysis is a standard test in a fabrication area. The test is made after any process change, equipment maintenance, or cleaning that could have the potential of contaminating the wafers. The C/V plot provides a wealth of other information, such as flat-band voltage and surface states. The thickness of the gate oxide can also be calculated from the C/V test data.

Contactless C/V Measurement

C/V monitoring as described above requires rigorous test-wafer preparation, which is time-consuming and expensive. Another method to determine unwanted charges (drift) and the other MOS gate parameters is with a contactless method called COS ([Fig. 14.45](#)). The MOS transistor method requires two electrodes separated by the gate oxide. A voltage on the top electrode (the metal) produces a charge buildup at the metal-oxide interface. A similar result can be obtained by using a corona source (the C of COS) to build a charge directly on the oxide surface. Increasing the charge puts the MOS transistor through the same paces as a metal electrode device and the same information, charge (drift), flat-band voltage, surface states and oxide thickness, can be calculated. The method correlates with the standard measurements.²¹

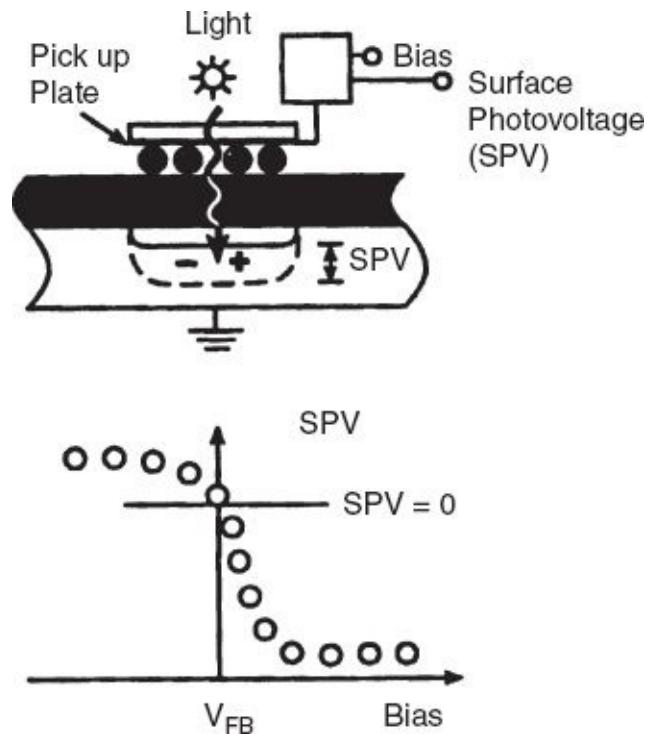


FIGURE 14.45 Noncontact surface charge measurement. (Source: Keithley Instruments Quantox™.)

Device Failure Analysis—Emission Microscopy

When a semiconductor device is operating, it gives off certain emissions of visible light. When there are certain problems, spots of visible light are given off at the trouble spot. For example, contamination-induced junction leakage will reveal itself as a light spot on the junction (at the surface). Microscopes fitted with sensitive detectors and charge-coupled imaging devices (CCDs) can locate and image these trouble spots. This method is particularly useful when electrical measurements indicate a failure in a sector of a circuit, but it cannot pinpoint the exact device(s) causing the failure.

Review Topics

Upon completion of this chapter, you should be able to:

1. Explain the difference between resistance, resistivity, and sheet resistance.
 2. Draw a sketch of the parts and current flow in a four-point probe.
 3. Compare the principles and uses of color interference, fringe counting, spectrophotometers, ellipsometers, and stylus for film-thickness measurements.
 4. Compare the principles and uses of groove and stain, SEM, and spreading resistance for junction depth measurements.
 5. List the methods and advantages of microscope and SEM inspection of wafer surfaces.
 6. Draw sketches of diodes in forward and reverse bias and their companion current-voltage curves.
 7. Explain the effect of surface-current leakage on a junction performance characteristic.
 8. Draw sketches of a bipolar and MOS transistor in operation and their companion current-voltage characteristics.
 9. List the process steps for a capacitance-voltage measurement and the principle of contamination detection.
 10. Describe the principle and use of atomic force microscopes.
-

References

1. McDonald, B., "Analytical Needs," *Semiconductor International*, Cahners Publishing, Jan. 1993:36.
2. Felch, S., "A Comparison of Three Techniques for Profiling Ultrathin p⁺-n Junctions," *Solid State Technology*, Jan. 1993:45.
3. Diobold, A. C. *In-Line Metrology, Handbook of Semiconductor Manufacturing Technology*, 2008, CRC, New York, NY:24–34.
4. Felch, S., "A Comparison of Three Techniques for Profiling Ultrathin p⁺-n Junctions," *Solid State Technology*, Jan. 1993:45.
5. Burggraaf, P., "Thin Film Metrology: Headed for a New Plateau," *Semiconductor International*, Cahners Publishing, Mar. 1994:57.
6. "Thin Film Measurements," Rudolph Engineering, product brochure, 1985.
7. Bordon, P., *NonDestructive USJ Characterization Using Carrier Illumination™ Measurements*, www.boxercross.com, May 2013.
8. McDonald, R., "How Will We Examine IC's in the Year 2000?" *Semiconductor International*, Cahners Publishing, Jan. 1994:46.
9. Braun, A., "Metrology Adapts to Meet CD Measurement Needs," *Semiconductor International*, Feb. 2002:73.
10. Braum, A., "Optical Microscope Continues to Meet High Resolution, Defect Detection Challenges," *Semiconductor International*, Dec. 1997:59.
11. Baliga, J., "Defect Detection on Patterned Wafers," *Semiconductor International*, May 1997:64.
12. *Product Description Bulletin*, Technical Instrument Company, San Francisco, CA, 1995.
13. Stapleton, J., *Optical Profilometry-Zygo NewView™ 7300*, Penn State Materials Characterization Laboratory, 2013.
14. Wolf, S. and Tauber, R., *Silicon Processing for the VLSI Era*, Lattice Press, Sunset Beach, CA, 1986:447.
15. KLA Tencor, *Candela CS20 Wafer Inspection System* product description, <http://www.kla-tencor.com/defect-inspection/candela-cs20.html>.
16. Burggraaf, P., "Patterned Wafer Inspection Now Required!" *Semiconductor International*, Cahners Publishing, Dec. 1994:57.
17. Peters, L., "AFMs: What Will Their Role Be?" *Semiconductor International*, Cahners Publishing, Aug. 1993:62.
18. "Time-of-Flight Secondary Ion Mass Spectrometer," brochure, Charles

Evans & Associates, Redwood City, CA, 1994.

[19.](#) Ibid.

[20.](#) Braum, A., "E-Beam Techniques Measures Product Wafer Composition, Thickness," *Semiconductor International*, Nov. 2001.

[21.](#) Peters, M., COS Testing Combines Expanded Charge Monitoring Capabilities with Reduced Costs, Keithley Instruments, Inc., product description paper.

22. Adams, T., "IC Failure Analysis: Using Real-Time Emission Microscopy," *Semiconductor International*, Cahners Publishing, Jul. 1993:148.

CHAPTER 15

The Business of Wafer Fabrication

Introduction

“The semiconductor story is unparalleled in the history of mankind.”

Dan Hutcheson, VLSI Research, Inc.

From its humble beginnings in laboratory-like production facilities to the automated wafer fabs of today, the semiconductor industry continues to evolve as the demands of new circuits, business factors, and global competition grow. While Moore's law still sets the product goalposts, fiscal considerations are setting the goals for production facilities, equipment, and processes. Overall productivity and cost of ownership are the factors that determine the bottom line. Strategies to increase the bottom line include maximizing equipment cost of ownership, automation, cost control, computer-automated manufacturing, computer-integrated manufacturing, and statistical process control.

Today the worldwide semiconductor industry has sales in the \$300 billion per year range and it is supported by a \$37 billion equipment industry and a \$47 billion materials industry.¹ Remarkably, it has stayed market-driven and spawned new products, from mainframe computers to desk and laptop computers, to the explosion of the hand-held wireless devices in the industry. This industry has created new worlds with cultural, economic, and international impacts unimaginable a few decades ago.

Driving these innovations have been the creative genius of circuit designers and the revolutionary leaps in manufacturing technology. Of these two things, it is manufacturing (microchip fabrication) that has produced products of increased performance, in ever-increasing volumes and at higher quality levels, and lower prices than ever before. This chapter examines these driving factors of microchip fabrication.

Moore's Law and the New Wafer-Fabrication Business

The semiconductor industry started supplying commercial products in the late 1940s. The manufacturing lines were little more than laboratories, and the workers were primarily trained technologists. In 1965, Gordon Moore at Intel Corporation developed his observation that transistor density was doubling every 18 to 24 months. Known as Moore's law, this observation not only predicted the future of the industry but became the industry's guide. It also set off an industry cycle that has driven the industry's growth.

The general implementation of Moore's law has been feature-size reduction (*scaling*) that leads to better performance and cost reduction, and in turn grows

markets and attracts investment that leads to the next uptick in performance ([Fig. 15.1](#)).

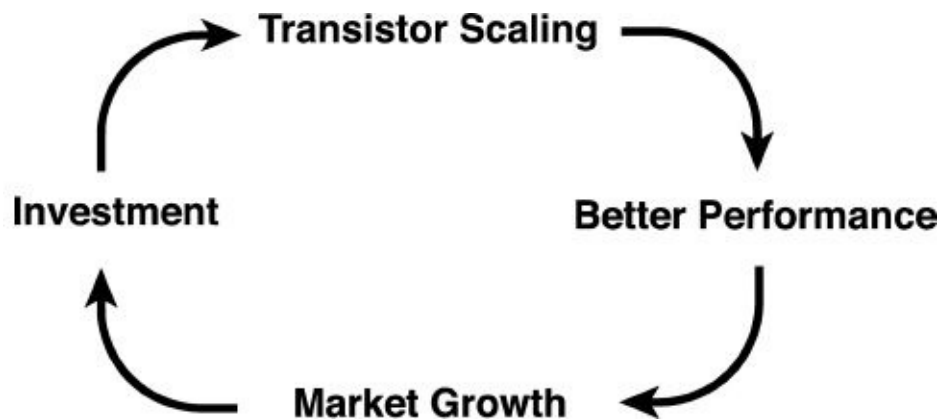


FIGURE 15.1 Moore's law-driven semiconductor business cycle.

By the 1970s, the manufacturing area had changed to cleanrooms, with highly specialized equipment attended by skilled production workers. A fabrication area of 2000 to 3000 ft² could be built at that time for \$2 to \$3 million.²

The VLSI/ULSI era saw wafer diameters grow to 200 mm, a wafer size too heavy and too valuable to process manually. Hence, the industry changed over to automation with Class 1 cleanrooms, and even more specialized and automated equipment and process-control systems. By the mid-2000s, the industry was shipping large microprocessor devices and entire systems combined on a chip (SoC). The industry became dominated by large companies as the cost of wafer-fabrication facilities rose to the \$8 to \$10 billion range (or gigadollars). The price tag is not so surprising, considering the factors driving the processes.³

Wafer-Fabrication Costs

A number of factors contribute to the cost of producing a functioning die ([Fig. 15.2](#)). They are generally divided into fixed and variable categories. Fixed costs are those that exist regardless of whether any die are being made or shipped. Variable costs are those that go up or down with the volume of product being produced.

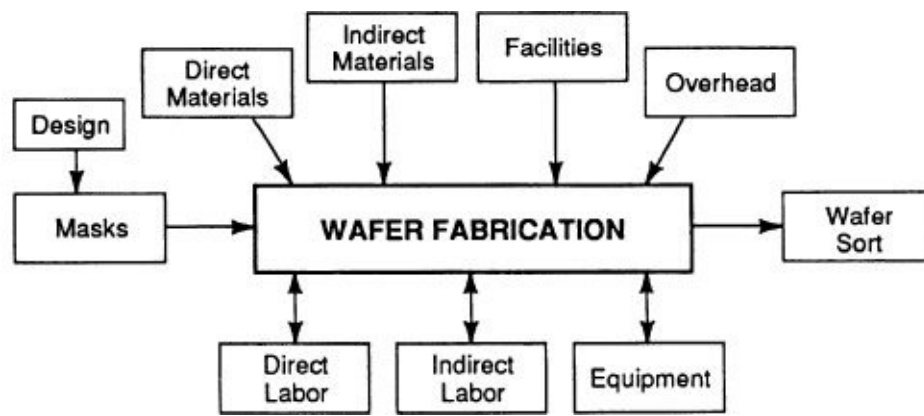


FIGURE 15.2 Fabrication cost factors.

Fixed costs include:

- Overhead—administration, facilities, and research
- Equipment

Variable costs include:

- Materials—direct and indirect
- Labor
- Yield

Overhead

Overhead costs are all those incurred by the administrative and executive staff plus the cost of providing and maintaining a facility. A curious fact of company growth is that beyond a certain level, the number of administrative personnel grows faster than the manufacturing workers. As companies grow, more information is generated internally, and more information must be handled from customers and suppliers.

To be effective, the information must be available to an ever-growing staff. The two needs result in more and more staff processing “information” rather than product. Also, decision making becomes more formalized (and costly) as more departments become stakeholders in the outcomes. For example, some 50 percent of the current workforce of the industrialized economies is involved in information processing.⁴ A primary overhead expense of semiconductor manufacturing is the design activity. With expensive CAD systems and a large professional design team, the cost of circuit design is considerable.

The cost and maintenance of the facility are major contributing costs. A fabrication area occupies only about 20 percent of a facility’s total area, yet creates the majority of the expense. Air conditioning, chemical storage and delivery, and the cost of the cleanroom are also major expenses. Fabrication cleanroom costs for

a ULSI facility are thousands of dollars per square foot. Actual floor costs are a factor of the cleanliness strategy chosen (see [Chap. 4](#)). A total cleanroom layout is more expensive overall than a hybrid or mini-environment approach. But in the latter, there is more expense for the equipment.

Many semiconductor companies maintain their own wafer-fabrication facilities and enjoy the benefits of controlling all the operations from design to packaging. They are known as *integrated device manufacturers* (IDMs). But the costs are high and an in-house fab will often be idle, thus driving up the overall fabrication costs. Two results of high-fab costs are the *fabless semiconductor company* and the *merchant foundry*. A fabless semiconductor company makes the circuit designs and contracts with a merchant foundry to actually perform the wafer-fabrication process. Packaging might be done at the foundry, or the wafer/chips may be moved to merchant packaging firm for this phase.

Materials

Manufacturing materials are divided into the categories of direct and indirect. Direct materials are those that go directly in or on the chip. This includes the wafers, the materials, and chemicals needed to form the deposited layers and doped layers, and the packaging material costs. Indirect materials are the masks and reticles, chemicals, stationery supplies, and other materials that support the process but do not enter into the product.

Equipment

This cost is the equipment used directly in the fabrication of the devices and wafers. It shows up in the cost calculations as fixed overhead or as depreciation. Depreciation is the loss of value of the machines as they wear out or become obsolete.

The transition to 300-mm wafers, and now 450-mm wafers, has bumped up equipment costs. At the transfer stage, 300-mm processing is essentially the same as 200-mm processing. This means that the process tools are the same, except for size capability. The larger wafers generally require more handling time and more process time. The overall effect is lower equipment cost per die ([Fig. 15.3](#)). These losses equate to more tools to maintain production quotas, and more expense. [Figure 15.3](#) compares the productivity of the primary process tools for the two diameters. New to the 300-mm process is a greater need for in-line metrology measurement and monitoring. The downside of larger wafer diameters is bigger losses if wafers are misprocessed or have low yields. Now, 300-mm processing is proceeding with copper metallization, which brings surface factors and new low-k

materials into the picture. Monitoring and controlling each of these steps at the process tool is critical if the whole system is to work at the end of the process. Critical dimension measurements and e-beam defect inspection systems have become in-process requirements and add to the equipment cost.

Tool	200-mm baseline (1.0x)	300-mm productivity in 2001	Expected productivity in 2002, 2003	Major upgrade required
Photo	1.0x @ wafer passes/day	0.6x	0.9x	Yes
Etch	1.0x @ wph	1.0x	1.2x	No
Wet bench	1.0x @ wph	0.8x	0.9x	No
CVD	1.0x @ preventive maintenance	0.3x	1.0x	No
PVD	1.0x @ available time	0.9x	1.0x	No
CMP	1.0x @ available time	0.8x	1.0x	No
Furnace	1.0x @ wafers/day	0.7x	0.9x	No
Implant	1.0x @ wafers/day	1.0x	1.0x	No
Defect inspection	1.0x @ wph	0.6x	1.0x	Yes
Metrology	1.0x @ wph	0.7x	1.0x	Yes

FIGURE 15.3 Comparison of 200- and 300-mm factors. (From Thomas Sonderman, Reaping the Benefits of the 450-mm Transition, *Semicon West*, 2011.)

Interestingly, the move to larger wafers and more sophisticated processes has increased the typical life cycle of new processes. [Figure 15.4](#) shows the history and projections of wafer-size life cycles.

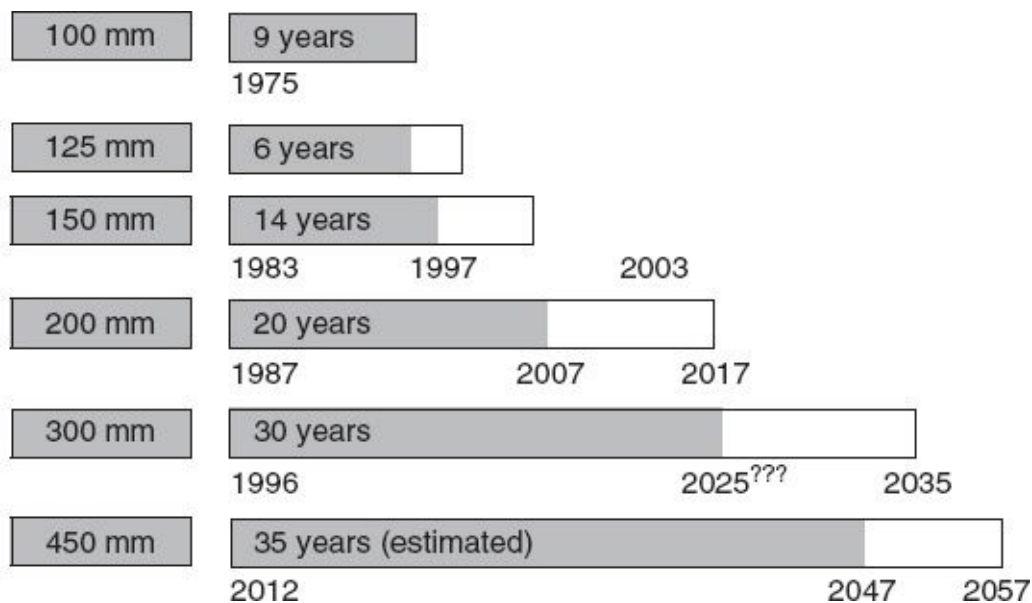


FIGURE 15.4 Wafer-size life cycles. (Courtesy of Future Fab International with 450-mm Update.)

Labor

Labor has both direct and indirect components. Direct labor takes in those workers actually handling the wafers and equipment. Indirect labor refers to support personnel such as supervisors, engineers, facility technicians, and administrative workers. Ironically, the new demands of process accuracy and productivity, using very sophisticated process tools, have returned the industry to the requirement that operators now have more technical training. Some companies, such as Intel, require all of their fabrication workers to have a technician-level education. All the preceding, coupled with higher numbers of technicians and engineers needed to set up and maintain the tools, have resulted in increasing labor costs over time.

Production Cost Factors

Example contributions of the factors to the overall unyielded die projected for 300-mm diameter wafer costs are shown in the following table. The percentages will vary with the wafer and feature size, degree of automation, and number of process steps. Depreciation in the chip industry is high as a result of the short useful lifetime of the expensive equipment. Direct materials include wafers. Others include the overhead costs.

Factor	Percent Contribution	Costs (\$)/wafer
Depreciation	35	1189
Labor	7	232
Maintenance	7	232
<u>Consumables:</u>		
Direct materials	12	405
Test wafers	6	203
Indirect materials	26	890
Other	7	226
Total	100	\$3378* (Example)

In any discussion about wafer cost factors, one is naturally curious about the actual costs. This number changes from line to line, with the complexity of the devices, and with the position of the product in the maturity cycle ([Fig. 15.5](#)). The more mature the product and process, the higher the yields and the lower the equipment-depreciation factor. Newer products suffer from operator learning

curves, equipment shakedown times, and development of new processes. Wafer volume and type are major influences on cost. High-volume products, such as dynamic random access memory (DRAM) generally have the lowest per-transistor or per-wafer cost because of high-manufacturing efficiencies. Application-specific integrated circuit (ASIC) wafers, on the other hand, have higher costs as a result of smaller production runs and the higher design and processing costs required for a varied product mix.

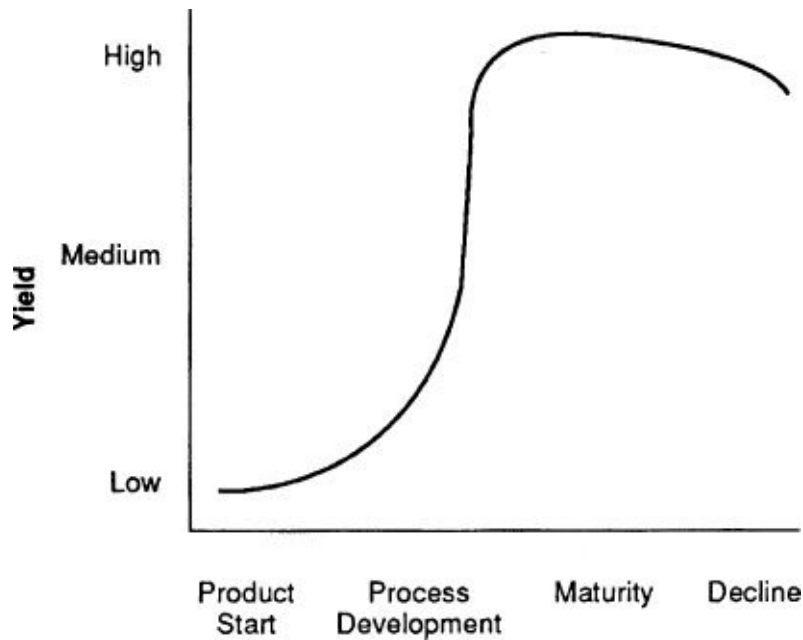


FIGURE 15.5 Wafer-production yield versus product maturity.

Yield

The overall fabrication yield (see [Chap. 6](#)) determines how the various costs affect the final die cost. If the die yield is low, the cost per die goes up. Not only are the fixed costs distributed over fewer die, but the variable costs go up as more materials are required to get out the die. When the yield is calculated into the cost, the term used is *yielded die cost*. The cost of producing the wafers without considering die yield is the *unyielded die cost*.

Die costs are a function of the wafer size, die size, and wafer-sort yield. For example, a \$3000 wafer-manufacturing cost for a wafer with 300 die will translate into an unyielded die cost of \$10 each. If the die sort yield is 50 percent, the die cost rises to \$20 each (\$3000 wafer cost divided by 150 functioning die). Increasing the sort yield to 90 percent would reduce the die cost to \$11.11.

Market pressures require that wafer-fabrication operations reach wafer-sort yields in the low 90 percent range, at ever faster rates. [Figure 15.6](#) shows some historic yield ramps for different levels of DRAM memory devices as they progress

from R&D to full production.

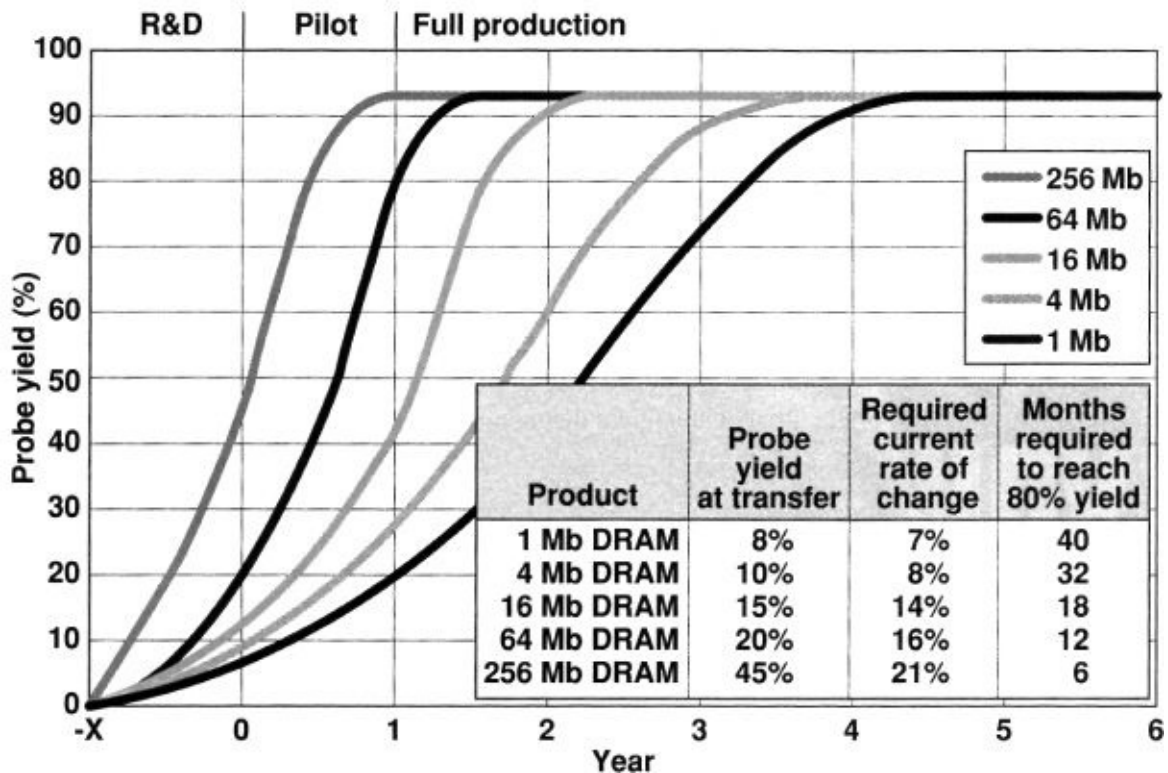


FIGURE 15.6 Probe yields for DRAM circuits. (Courtesy of Semiconductor International, January 1998.)

Yield Improvements

While it is true that increased attention is being given to traditional business factors, no less attention is paid to process and wafer yields. The effect of a yield improvement can be significant in terms of dollars. Consider the yield figures in [Fig. 15.7](#). Line 2 shows that a wafer-fabrication yield improvement of five percentage points increases the overall yield from 38 to 40.4 percent. If the fabrication area starts 10,000 wafers per month and there are 350 die per wafer and the selling price is \$5 per die, the increased revenue is:

Process	Change	Yield		Assembly	Overall
		Fab	Sort		
Base Line		0.80	0.50	0.95	0.38 (38%)
	Fab to 0.81	0.81	0.50	0.95	0.384 (38.4%)
	Fab to 0.85	0.85	0.50	0.95	0.404 (40.4%)

FIGURE 15.7 Yield impact of fabrication-yield improvement.

$$10,000 \text{ wafer starts} \times 350 \text{ die per wafer} \times 0.012 (1.2\%) = \$210,000/\text{month}$$

Whether this amount of increased income is significant depends on the cost of effecting the improvement.

Yield and Productivity

Yield has been the traditional measure of wafer-fab success. High yields have translated to lower production costs and higher profits. With wafer-sort yields at the 90 percent level, the next cost-reduction factor is productivity. Improved production efficiency is measured by the number of wafers per square foot of the total fab area (or number of wafers per piece of equipment), or the number of wafers produced per operator. Increasing these factors means a lower manufacturing cost. Throughput is another factor. Increasing the number of wafers out per hour means a more efficient process with lower costs. These factors increase the bottom line if the cost of achieving them is less than the advantage.

Increasing Wafer Diameters

Circuit functionality and speed are accomplished as dimensions move to the nano-level (one-billionth of a meter). However, this results in smaller components crowding the chip surface, forcing additional layers of metal connectors above the chip surface, and more process steps. Smaller component sizes on the chip surface require shallower doping layers and thinner dielectric layers. These requirements also increase the number of process steps. Along with higher-density chips, capacity improvements lead to larger chip size. However, larger chips on the same diameter wafer reduce productivity due to the lower chip per wafer count. The natural solution to this problem is larger-diameter wafers. Thus the industry has consistently increased wafer diameters as chip density has increased. The leading-edge production-level wafers now are 300-mm wafers (approximately 12 in), and an upgrade to 450-mm wafers (approximately 18 in) is underway. Like most volume-based manufacturing processes, increased volume brings down the price over time ([Fig. 15.8](#)).

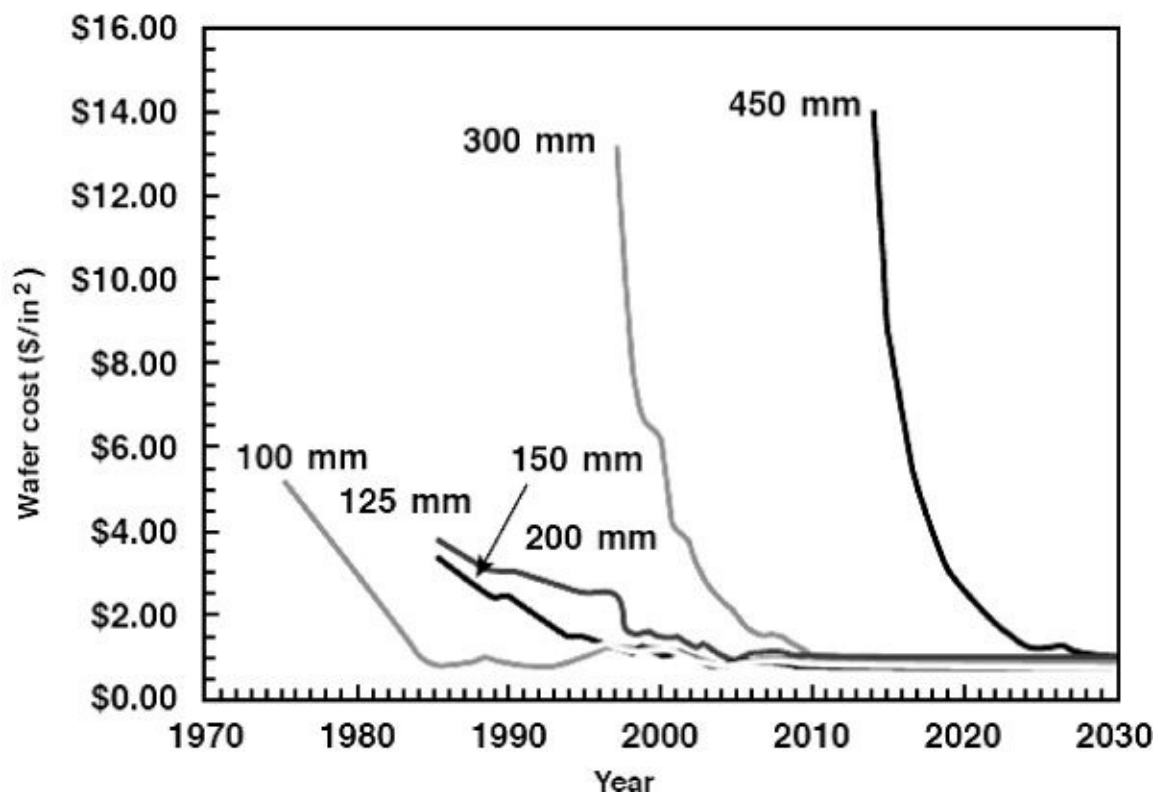


FIGURE 15.8 Wafer costs over time. (Scotten W. Jones 3200, ICK Knowledge LLC, www.icknowledge.com.)

The larger-diameter wafers have brought new production challenges and forced higher levels of automation. The breakpoint requiring automation was 200-mm wafer diameters. At a practical level, the weight and value of a 25-wafer batch of 200-mm wafers in a carrier became too high for manual processing. Entegris Wafer Handling calculates a 25-wafer batch of 460-mm wafers weighs in at 19 pounds, without the carrier. Maintaining productivity and yields with larger wafers creates higher costs. The equipment to process larger wafers with tighter tolerances becomes more expensive. Improvements in processes and equipment are needed to maintain uniformity across a larger wafer (or more tools processing single wafers). Of course, all of the materials and processing environment must become cleaner, since the smaller, more closely packed devices are more sensitive. And don't forget that at these levels, the number of tests and characterizations increases to maintain product quality and process control in a wafer-fabrication line that is rapidly moving large volumes of wafers. Of special note is the increase in process steps at the ULSI/nano level. Squeezing the devices closer together and making them smaller has introduced problems that were solved by the addition of new process steps, such as planarizing techniques to overcome topography-generated imaging problems. The additional steps drive up the process and inventory expense.

Costs in moving from 300-mm wafers to 450-mm wafers are projected to rise due to 30 to 100 percent larger fab areas, 20 to 50 percent increase in equipment costs, 10 to 30 percent increase in "beam" tools (longer process times), and a 1.7

percent increase in consumables. Yet overall the transition is projected to net a 25 percent decrease each in capital expenditures and process costs at the 22-nm node.

The move to 300-mm wafers brought the option of smaller fabs (or *mini-fabs*) running 10,000 wafers per month and producing the same number of chips as a 200-mm fab.⁵ Similar advantages are available with 450-mm wafers.

Microchip fabrication costs and management are further complicated by the uncertainties that come from free-market pressures and the rapid turnover of technologies. A life cycle of a product usually starts with rapid growth that may require new equipment, processes, and/or an upgraded facility. The volatility may create supply imbalances such as silicon or silicon wafer shortages. Lastly, a product cycle often experiences price fluctuations as competitors catch up or move the technology to the next level. All of these occur in a market where sales prices erode over time.

“For example, the decline in price per bit has been stunning. In 1954, five years before the integrated circuit was invented, the average selling price of a transistor was \$5.52. Fifty years later, in 2004, this had dropped to a billionth of a dollar. A year later in 2005 the cost per bit of dynamic random access memory (DRAM) is an astounding one nanodollar (one billionth of a dollar).”⁶

Despite all the changes in processes, costs, and markets, the overall financial measurement of a fabrication area has remained the same: it is *the cost per functioning die shipped out of fabrication*. When extended to a complete merchant facility with assembly capabilities, the measure becomes *the cost per die shipped*. In the world of the megachips, the cost per transistor is becoming an indicator parameter.⁷ These measurements apply to in-house wafer-fabrication operations, foundry operations, and fabless chip companies.

Book-to-Bill Ratio

On the business side, one of the most-watched industry factors is the book-to-bill ratio (b/b). *Book* refers to bookings, which is the dollar amount of sales orders received in a period. *Bill* refers to the dollar amount of billing invoices sent out. In a normal operation, chip orders are shipped and billed some time after the order is received. In good times, the delay reaches into months or longer if shortages occur. A high book-to-bill means that unfilled orders are accumulating, generally indicating a healthy economy. A low book-to-bill means that orders are drying up, and the factory is shipping out to inventory rather than building new product. This indicator is used to project inventory and production levels. A b/b ratio greater than 1 indicates a rising market, and less than 1 indicates a declining market. When the market becomes very good, a high b/b ratio is worrisome, since shipping times become stretched, and the end user can become starved for chips.

Cost of Ownership

Shifting tool or equipment purchase decisions away from purely technical to business-related factors has driven the development of *cost of ownership* (CoO) models. These models attempt to bring together all of the relevant factors driving the total cost of ownership of a tool, process, or facility over its expected lifetime. Beyond the initial equipment purchase price, tools differ in the amount of expensive floor space they occupy, the amount of power and materials they require for operation, the yield of in-spec wafers, maintenance, repair and failure rates, and so forth. The CoO formula is one developed by Sematech to evaluate equipment purchases.⁸

$$C_w = \frac{\$F + \$V + \$Y}{L \times TPT \times Y_{TPT} \times U}$$

where C_w = cost of the finished wafer, based on the particular cost associated with the tool or process for the lifetime of the tool

$\$F$ = fixed costs, which include the equipment purchase prices, facility costs, initial modifications, and so forth

$\$V$ = all material, labor, and process costs generated when the tool or process is operating

$\$Y$ = the cost of wafers scrapped due to tool scrap and defect-induced losses

L = the lifetime of the tool in hours

the wafer throughput rate as reduced from the maximum by

TPT = maintenance requirements, setup, test wafer monitoring, and so on, expressed in wafers per hour

Y_{TPT} = yield factor

U = the tool utilization factors that reduce available process time from the maximum

Each of these equation terms is calculated from formulas that consider the subfactors for a particular process. The CoO formula provides the method to determine tradeoffs for various tool factors. For example, a tool may have a high initial cost that is offset by lowered operating costs or a long time between failures. Or, a tool may provide a high yield but require so much adjustment and calibration that additional machines would be required to meet production schedules.

Automation

As most industries mature, the technology becomes stabilized, and the market drives up demand. These two conditions are precursors to process automation. Since 1940, automation of oil refineries has reduced the number of workers by a factor of 5.⁹ Automation of semiconductor processes has been in progress since the 1970s, when process tools were designed to accept wafers in cassettes. Since then, automation has marched along toward the fabled dream of the total “peopleless lights-out” fab. Automation stages start with the process tool and extend to the factory level.

Process Automation

The first level of automation is of the process itself. Most semiconductor equipment by definition, automates a part of a process. Photoresist spinners automatically dispense the primer and resist at the correct speeds and for the correct time. Automatic gas-flow controllers dispense gases to the tool in the right amount, at the right pressure, and for the right time. Process automation brings consistency to the process and the product by reducing reliance on operator skills, training, morale, and fatigue.

Most tools are controlled by a set of instructions programmed in an onboard computer. The program is called a *recipe*. The recipe is loaded into the machine by the operator or from a central host computer.

Wafer-Loading Automation

The next level of automation is the loading and unloading of the wafers. The industry has settled on the *front-opening unified pod* (FOUP) as the primary wafer holder and transfer vehicle. FOUPs are placed on the machine by various mechanisms. Elevators and/or wafer extractors or robots feed the wafers into the particular process chamber, spin chuck, and so on. In some processes such as process tubes, the entire cassette is placed in the process chamber. This level of automation is referred to as the “one-button” operation. With one button, the operator activates the loading system, and the wafers are processed and returned to the cassette. At the end of the cycle, the machine sounds an alarm or turns on a light, and the operator removes the cassette.

Some machines have buffer storage systems that maximize the machine efficiency by always having fresh wafers (or reticles, for the imaging tools) available for processing. These are called *stockers*. The operator places the FOUP(s) on the machine loader and pushes a start button, after which the machine takes over the processing. At 300-mm diameter and above wafer levels, and with single-wafer processing tools, wafers may be transported in individual holders.

Clustering

Mating two or more process steps in a single unit is another level of automation. Generally, this level is called *clustering*. The industry has been “clustering” for a long time. For example, photoresist spinners were long ago mated to soft-bake modules and other track equipment groupings.

Recent clustering designs ([Fig. 15.9](#)) have been driven by both technical and economic forces. On the technical side, some clustered processes make better products by, for example, keeping a silicon wafer clean after etch and before metal deposition. Another process advantage is sequential deposition of different materials in the same chamber. In these cases, the deposition process is better, because the wafer is not exposed to air between steps. Any time a wafer loading or unloading step can be eliminated, both cleanliness and cost are favorably affected. For vacuum processes, time and cleanliness factors are affected when two or more processes can be performed with only one vacuum pump down. Clustering in which two or more sequential processes are performed is called *integrated processing*.¹⁰ On the economic side, some types of processes that are clustered for increased throughput are called *parallel processing*.

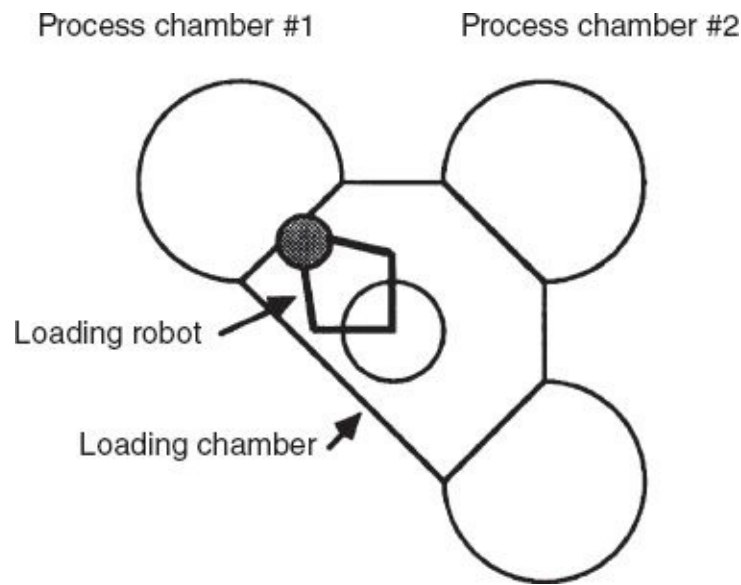


FIGURE 15.9 Three-chamber cluster tool.

Despite the obvious attractions of clustering, there are drawbacks. Clustering for critical process advantages is easy to justify. But a cluster of same-type processes requires interlocks, electronics, and software more sophisticated than the individual tools. And a shutdown for maintenance or repair idles a larger part of the production capacity. Another barrier occurs when cluster modules cannot be provided by the same vendor. Customers prefer one responsible vendor, and vendors are slow to pair up when responsibilities are clouded.

Wafer-Delivery Automation

The third level of automation is when wafers are automatically brought to, loaded on, and removed from the machines. In addition to the production advantages of automation, there are ergonomic, safety, and cost benefits. These accrue from the weight of a batch of larger-diameter wafers in a FOUP, which can reach the 18 to 20 lb (6.7 to 7.5 kg) level. Operator injuries are a possibility, and the financial loss of dropping a batch of expensive wafers can be staggering. Early delivery systems used traveling robotic carts that duplicated human delivery ([Fig. 15.10](#)). Called *automated guided vehicles* (AGVs), the carts travel along the aisles and dispense wafer cassettes when the machines need them. Another version is the rail-guided vehicle (RGV), where carts follow a track laid out between the process tools and the stockers. These approaches have the advantage of being retrofittable to fabrication lines where the equipment is lined up in rows.

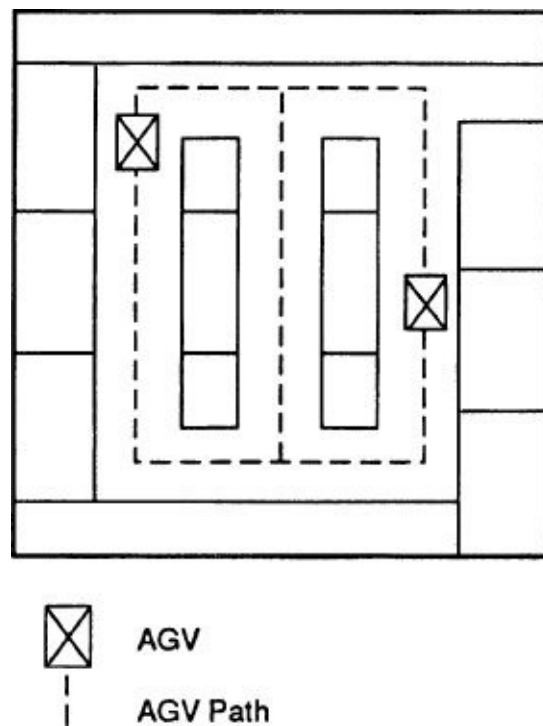


FIGURE 15.10 Automated guided vehicle wafer-delivery system.

The favored approach is the use of an overhead rail (also called a gantry). The wafer FOUPs arrive at the process tool area ([Fig. 15.11](#)), where a secondary system (usually a robot) removes them from the overhead rail and places them in the tool buffer. This system works best with machines that are grouped in bays rather than in the traditional linear layout.

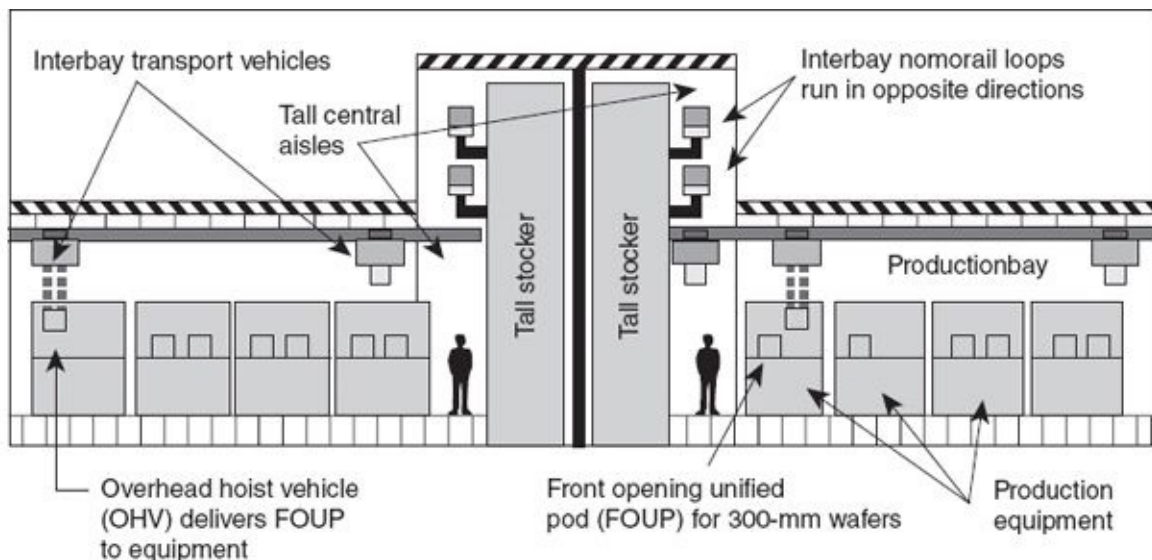


FIGURE 15.11 Overhead rail or stocker wafer-delivery system.

Closed-Loop Control-System Automation

The industry is just embarking on the final level of total automation, the closed-loop feedback system. There are two aspects. Some tools have onboard sensors that measure critical process parameters. Through feedback electronic circuits, these tools adjust themselves to maintain process-operation specifications. This stage of automation is difficult to achieve for many processes. It requires measuring sensors that operate in hostile environments, such as heated-tube furnaces and sputtering chambers. One area of progress is in the area of mask-making, where recognition systems are advanced enough to compare the manufactured plates with the design criteria.

The second aspect of this final level of automation is tools that have, or are connected to, automatic inspection or measuring subsystems. The subsystems measure important parameters in real time (i.e., as they are happening), compare the measurements with a standard, and feed the information back. A couple possible results of the feedback are that it may trigger a shutdown of an out-of-spec tool or make process parameter adjustments. The development of closed-loop process and machine control is necessary before the goal of a “peopleless, lights-out” wafer fab can be achieved.

Factory-Level Automation

The advent of higher levels of process and tool automation and inventory control systems requires higher levels of centralized control and information sharing. Most companies have computer-based management information systems (MISs) handling the paperwork and details of employment and finances. These systems are being expanded to the entire manufacturing environment in a process called *computer-integrated manufacturing* (CIM).

CIM is the computerization of all plant operations and the integration of those operations into one computer design, control, and distribution system. The processes involved in CIM are all related and interdependent, as illustrated in [Fig. 15.12](#). The major activities of CIM are business functions, product design (mask and circuit), manufacturing planning (inventory, shop floor priorities, etc.), manufacturing control, and the fabrication processes. A complete CIM system is interactive at all levels. This means that each of the five functional areas input data to the system in real time and that the information is available to all who need it.

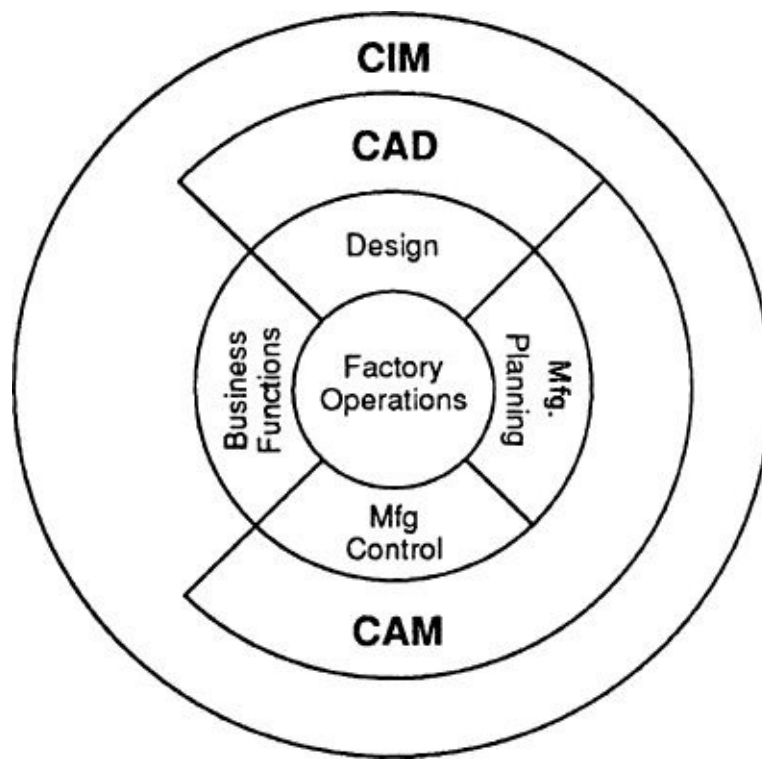


FIGURE 15.12 Provinces of factory computer control systems.

Computer-aided design (CAD) and computer-aided manufacturing (CAM) are two subsystems within the CIM system. The role of CAD has been discussed in terms of mask and circuit design. CAM is the part of the system that does the planning and control of the manufacturing operation. A CAM system includes a computer network and the automated process tools and material-delivery systems.

In concept, a CIM system kicks into operation when a customer order is received. The computer logs the order and initiates the CAD system to start the design (if it is a custom order). Logging of the order also triggers the ordering of needed materials in the right quantities and schedules their delivery times. This subsystem goes by the name *computer-aided process planning* (CAPP). Through the CAM program, the individual process recipes are downloaded to the individual equipment computers. Once processing begins, the CAM system controls the WIP (work in progress) and makes necessary priority decisions to meet shipping schedules. It also keeps track of equipment performance and schedules repairs and maintenance.

An important feature of the CAM system is yield monitoring and reporting. Important measuring systems are connected directly to the factory computer. If a poor yield problem appears, the system reports it to the engineering and facility staff, and (if necessary) will reorder materials to make up the losses. At the end of the production run, the system calculates costs and yields and schedules shipment to the customer.

Other packages included in CIM systems are facility monitoring, process modeling, and security systems. Facility monitoring might include power

consumption and environmental factors inside the plant. Some facility CAM systems include monitoring the levels of liquids and gases in storage with automatic alerting and/or reorder. Monitoring and precise control of these factors can save appreciable amounts of money. Process modeling is a system of testing a particular design against the known process variations. The computer can run many variations that simulate the changes expected in fabrication. Good modeling can identify weak points in the design or process before wafers are committed to the line. Security monitoring may include the entering and exiting of employees (and/or intruders) and the securing of expensive finished products or materials. A security system may also include fire and other hazard controls.

Equipment Standards

Given the many suppliers of materials and equipment and the very stringent technical demands of chip manufacturers, the need for standards is obvious. In many industries, different manufacturers attempt to establish their “standard” as the industry standard, as a way of maintaining a competitive edge. Fortunately, the semiconductor industry informally has settled on many standards—for example, the use of standardized cassettes.

In 1973, Semiconductor Equipment and Materials International (SEMI) established a standards program. Supplier and user personnel come together and establish standards by a consensus process. At the automation level, SEMI has published communication protocols for equipment interfaces.

Fab Floor Layout

The tradition fab floor layout of lining up the equipment in rows has given way to the bay (or tunnel) system. This system groups sections of the process in separate rooms, which minimizes contamination from personnel traffic and cross-contamination from other process areas. Automation needs are forcing a rethinking of the most efficient way to move material while maintaining cleanliness.

Robots and overhead gantries have become the material movers of choice.¹¹ Also driving robot use is the increasing weight of wafer lots (larger diameters), especially in FOUP-type transfer modules.¹² Robotic design and performance challenges come to the forefront in vacuum tools, especially when the process uses corrosive gases.

One equipment layout option is the process island. The various tools for a particular process segment are grouped around a single loading or unloading robot. The tools may be single or in clusters as shown in [Fig. 15.13](#).

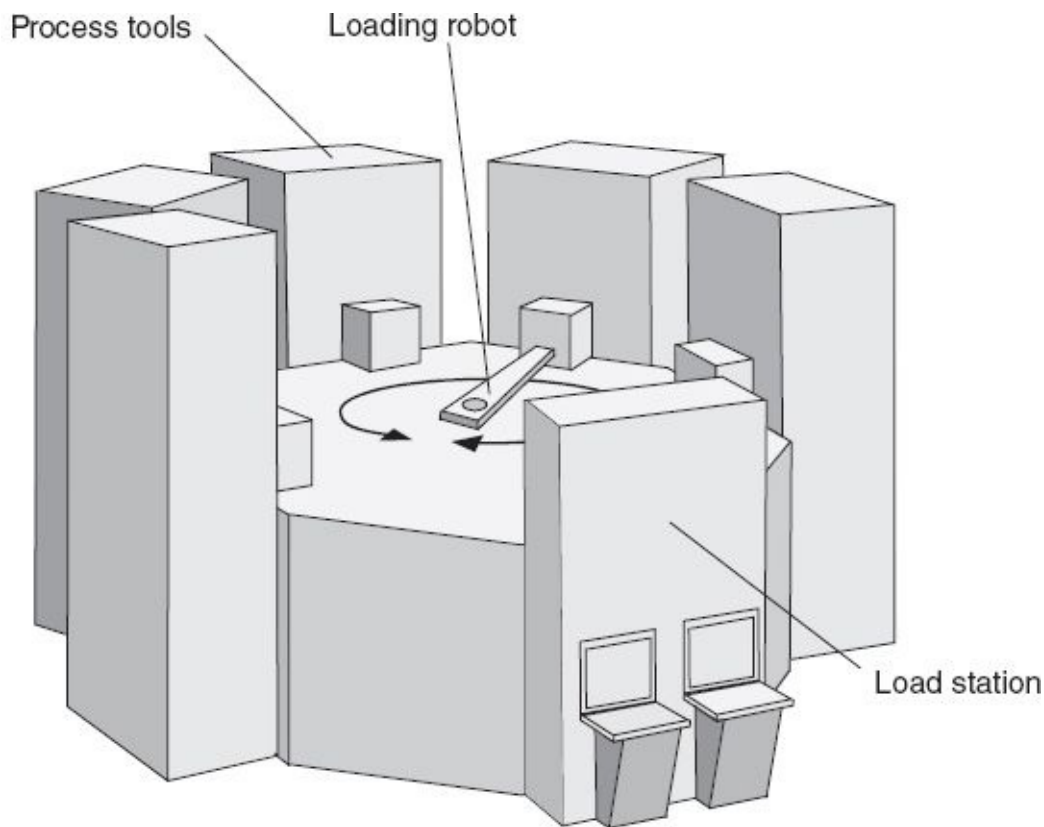


FIGURE 15.13 Example etcher process cluster for 450-mm wafers.

Batch versus Single-Wafer Processing

For the first two decades of process development, the drive was to ever-larger batches of wafers. Simple economies of scale were behind the drive, that is, gaining more product output per process sequence. However, the arrival of larger-diameter wafers and VLSI/ULSI-level circuits has changed the picture. Larger wafers require larger process chambers, which in turn push the limits of uniformity. Many processes, particularly the deposition steps, are easier to control in a smaller (single-wafer) chamber. Fast-cycle time fabs (ASIC and small volume devices) can use the flexibility of single-wafer processing to good advantage.¹³

Offsetting-improved process control is a loss of productivity as compared to batch processing. Larger-diameter wafers help somewhat (300-mm diameter wafers have almost four times the area of 150-mm wafers). Single-wafer process productivity is improved by tool designs such as a design that allows one vacuum pump to be down for an entire batch while wafers are processed individually, and other improvements.

On the downside, single-wafer processing requires higher levels of repeatability. Assuming the same-size wafers, a single-wafer processor will have to perform 25 times compared to a one-batch 25-wafer batch process. The equipment for multiple-batch processes is also more expensive than equipment for single-batch processes. Each single-wafer processor requires essentially the same components as a batch

tool.

Green Fabs

Additional pressure on wafer-fabrication processes is generated by environmental laws and environmental concerns. A wafer-fabrication operation uses many hazardous chemicals and produces waste products. In-plant controls of chemical handling, storage, and use are a part of fab *environmental, safety, and health* (ESH) programs. Active research is aimed at the development of processes that use less chemistry and chemicals that are less harmful. Also, the SIA roadmap projects water consumption dropping from 30 gal/wafer to 2 gal/wafer.¹⁴ The roadmap also calls for lowered energy use per wafer, and the elimination of polyvinyl chlorides (PVCs). Along with cost savings due to the preceding changes and safety improvements, there are cost savings related to lower-hazard chemicals in terms of the costs of onsite treatment and storage, transportation, and acceptance at special waste facilities.

Statistical Process Control

This text for the most part has presented the process technology with few mathematical formulas. The next few sections on statistical process control will break from that approach. Statistical process control (SPC) is a powerful and needed tool to maintain process control and improve yields. Unfortunately, SPC is based on statistics, which are applied with the language of mathematics. We will attempt to convey the background, use, and interpretation of commonly used statistical tools with illustrative charts and graphs. No formulas will be presented.

The first question to address is, “What is the object of process control?” The answer is simple: to produce a product that falls within its design and operational specifications, and to do it at a high-enough rate to be profitable. Challenging this simple goal are processes that are influenced by a multitude of parameters, including the conditions already on the wafer.

Processes do not run in control without monitoring and adjustments. Process control provides the information to make the necessary changes. Statistical process control does it in a way that relies on established mathematical principles that govern the type of processes used to make chips. Process-control techniques can vary from simple to very complex. The simplest (and most familiar) is the calculation of the average of a group of numbers (called a *population*). We also know that the extremes of a population (range), along with the average, give us an idea of the distribution of the data within the population.

For example, the two groups in columns A and B of [Fig. 15.14](#) have the same

average. If the numbers were the sheet resistances of wafers from furnace A and B, we could easily come to the decision that furnace A has the most in-control process because of its tighter distribution. This fact is illustrated when the data is plotted ([Fig. 15.15](#)). The plots are called *histograms* and visually display data distributions that a simple average calculation will not reveal. Average calculations and histograms are statistics in action. Histograms are usually the first step in determining whether a process is in control.

Reading no.	Run A	Run B
1.	25	28
2.	24	25
3.	26	23
4.	23	26
5.	25	25
6.	26	22
7.	25	23
8.	24	25
9.	25	27
10.	26	25
Average	24.9	24.9
Range	26 - 23 = 3	28 - 22 = 6

FIGURE 15.14 Sheet resistance reading ($\Omega/\square$).

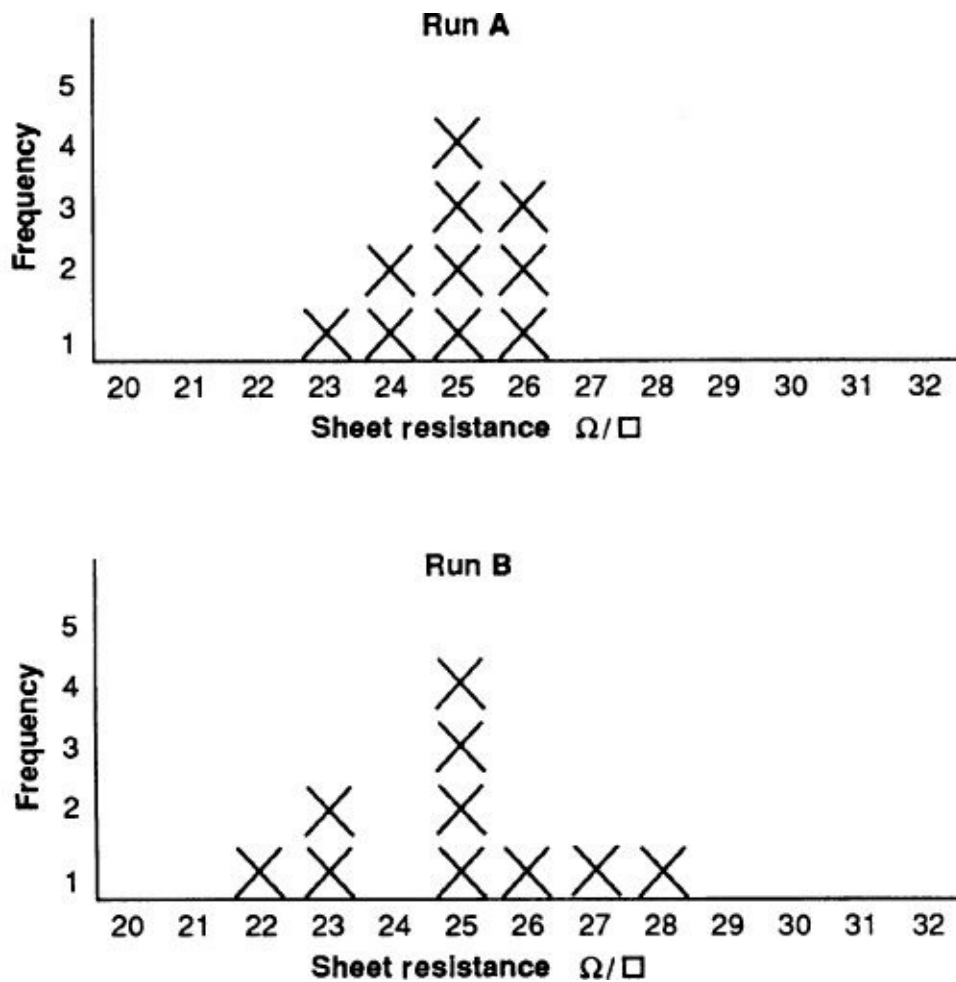


FIGURE 15.15 Frequency distribution of sheet resistances.

Their power comes from a mathematical distribution known as the Gaussian distribution. Named after the famed mathematician Karl Friedrich Gauss (1777 to 1855), its origin is interesting. Gauss set out to reconcile the different star positions reported by different astronomers. His approach was to make all the necessary corrections to the observations, taking into account the time of the year and from where on the Earth the observations were made. He expected that, when all the corrections were made, there would be agreement from all the position calculations on the position of a particular star. After all, reason dictates that a star can occupy only one position at a time, and we should be able to determine that position.

However, the final data did not confirm his hypothesis. After all corrections were made, there were still a number of locations for each star. Fortunately, Gauss did not scrap the project but went on to establish the basis for the field of statistics and distribution probabilities. If the math-averse readers will hang on, they will find that the concept of probability is not so difficult. The way Gauss analyzed the data was to plot the various calculated locations for a star. He calculated the center point (average) and drew a circle that encompassed the star location the farthest from the center. He reasoned that there was a 100 percent probability (a probability of 1) that

the real star location was within that circle. He also reasoned that the probability of finding the star within the confines of smaller circles was less than 1. In fact, the smaller the circle, the lower the probability that the star was within those boundaries.

It also turns out that processes produce data distributed like the example given. The mathematical distribution is known as the Gaussian distribution. A good example is the height of blades of grass in a lawn.

If all the blade heights are measured and plotted on a histogram, the distribution will be the familiar bell-shaped curve, also called a *normal curve* ([Fig. 15.16](#)). From a probability consideration, it predicts that there is a higher probability that any given blade of grass will have a height closer to the average (center values) and that there is a lower probability that any given blade will be very short or very tall. Some other mathematical conditions that result in the distribution include human heights, IQ distributions, and (in most cases) semiconductor process parameter distributions, such as sheet resistance.

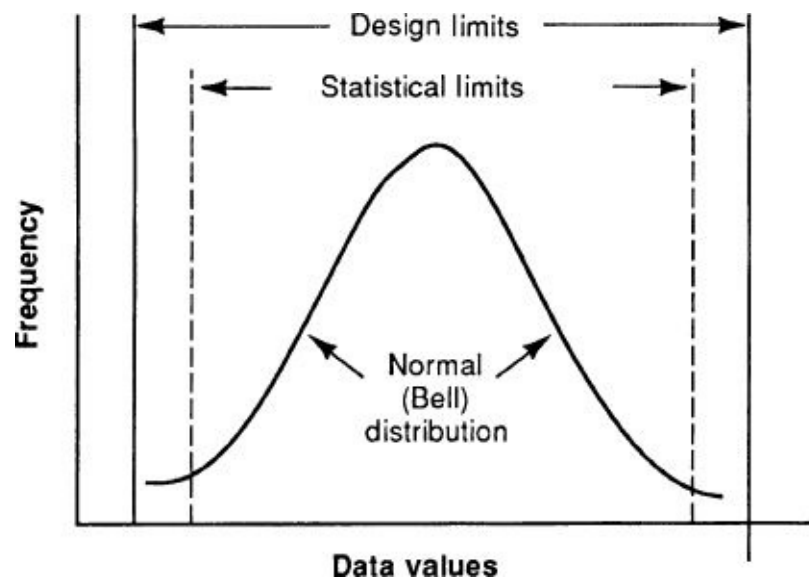


FIGURE 15.16 Normal distribution (bell) curve.

A first step in process control is to make a histogram of the particular process parameter and determine if the distribution is a normal distribution. If it is not, the chances are good that there is something wrong in the process. If the distribution is a normal one, the next step is to compare the range of the distribution with the design limits for the particular parameter ([Fig. 15.16](#)). This comparison is made to determine if the natural process distribution limits fall within the design limits. If they do not, the process must be fixed or some percentage of the parameter readings (and the wafers) will always be out of specification.

So far, the statistical methods explained are the after-the-fact histories of a

process. A most powerful statistic of real-time process control is the X - R control chart (Fig. 15.17). The chart is constructed in two plots. The lower plot shows the moving range R , which is descriptive of process stability and the parts, with the y axis representing the parameter values. The top graph has a horizontal line for the historic average of the parameter. On each side of the average are control limit lines that are calculated from the historic data.¹⁵ A control limit represents the limits that the individual data values will range between when the process is in control. Also on the chart are the process or design limits that represent the extremes the individual data points may have before being rejected. The bottom graph is constructed by calculating and plotting the amount that each data point varies from the average. When plotted, these values give further visual evidence of the amount of control in the process.

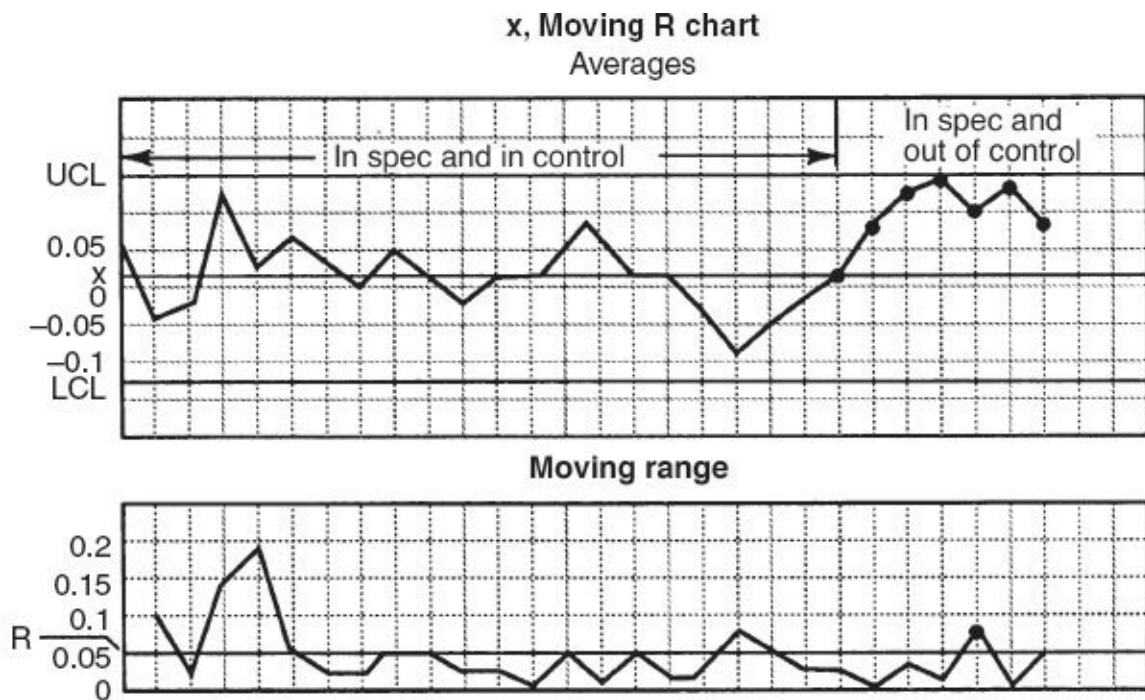


FIGURE 15.17 Moving R chart contains the averages of measurements, x , in the upper plot. The lower plot shows the moving range, R , which is descriptive of process stability.

The value of the X - R bar control charts is their predictive powers. A process in control will produce data points that tend to vary in a regular pattern about the average (top of Fig. 15.17). The mathematics of a controlled process predicts this regular fluctuation. It also predicts *when* a process is going out of control *before the data points exceed the control limits* (bottom of Fig. 15.17). The data points in part B have shifted to the top of the control-limit range. This is an unnatural pattern for an in-control process. When this situation occurs, the production operators, who maintain the charts as the data is produced, alert the proper personnel so the process

can be brought back into control *before the data points exceed the control or design limits* and wafers have to be scrapped. A number of more sophisticated controls used in processing are beyond the scope of this text.

Another powerful statistical tool is multivariable experiment analysis. Most measured quality-control parameters (sheet resistance, line width, junction depth, and so on) are influenced by a number of variables in the process. Line width, for example, varies with the resist solution, film thickness, exposure radiation time and intensity, baking temperatures, and etch factors. Any one or all can contribute to an out-of-spec condition. Multivariable evaluations allow the engineer to run tests that separate and identify the contribution(s) of each of the individual variables.

Designing an SPC system for a process requires selection of the proper statistical tool. Another decision revolves around the proper “indicator” population. Profit demands that all the die on all the wafers from all the batches, day in and day out, meet the specifications. However, picking the parameter population is not always an easy chore. Depending on the process, there are variations across the wafer, there are variations from wafer to wafer within a batch, and there are variations from tool to tool. Since every die cannot be measured, selecting the right sample point and sampling level becomes a demanding task.¹⁶

Inventory Control

A critical issue in fabrication cost control and yield is inventory level and control. As the number of processing steps has increased, so has the length of the processing time and the number of wafers in process (WIP). The problem is that the company pays for the wafers when purchased and doesn’t receive payment until the finished devices are shipped. This period can vary from two to eight months for a production line making similar circuits. Fabrication lines doing ASIC circuits have an even heavier burden when many different circuit types are going through the system. To get an idea of the burden, consider a CMOS-type process with 50 major steps and 4 substeps each, for a total of 200 processes. The high cost of the equipment generally requires some buffer inventory to ensure that the machines are operating at maximum efficiency. If each buffer has 4 FOUPs of 25 wafers each, the total inventory of WIP is 40,000 wafers. At a cost of \$100 per wafer (large diameter), the total inventory burden becomes \$4 million.

Excessive WIP affects productivity by having the capability of hiding process and equipment problems.¹⁷ With a lot of inventory at the stations, wafers can keep flowing out the back-end, while parts of the process are shut down. With a lower WIP, these problems are readily apparent and force the solution of problems. WIP also influences the overall fabrication yield. The collective experience of the industry is that the longer wafers are in the process, the lower their wafer sort yield.

Just-in-Time Inventory Control

Just-in-time inventory control is a philosophy based on the objective of, “Make only what is required, and only as required.”¹⁸ The system is simple in concept. All buffer inventories are reduced to an absolute minimum, from the storeroom to the machine buffers. To work effectively, excellent vendor relationships must be established. The incoming materials must be of the highest quality, as JIT leaves little cushion for extensive incoming inspections and returns. Second, the vendor is asked to maintain ready-to-ship materials at its facility. In effect, the vendors are asked to hold the inventory previously held by the chip manufacturer, a situation that they are not always happy about. The chip manufacturers have to be very efficient in assessing their chips’ quantity, quality, and delivery needs, and must have a system that expedites material to the proper process tool and detects quality problems quickly.

JIT also has applications in the process flow. Some companies will commit wafers to a process only if the work can flow unimpeded through all of the substeps. This system gives higher yields and, if on a properly balanced line, increases throughput time even though parts of the line are idle. This system is called *demand-pull*. Wafers “upstream” are worked on only when there is demand from downstream process stations about to run dry of wafers.

An effective JIT procedure can reduce the number of operators in the fabrication area, since a good proportion of their time is spent sorting, staging, and delivering work to the process tools. Fabrication-area layouts can change from the traditional linear arrangement of process stations. A linear layout benefits a line that makes only a few different product types. Wafers are introduced into the front of the line and physically move on a first in-first out (FIFO) basis. Problems arise when particular lots of wafers must be moved quickly. Usually tagged as *hot lots*, they are given priority processing at each station. Besides the control problems associated with keeping track of the lots, hot lots have the effect of making regular products, sit in a queue waiting for the hot lots to clear. This layout is particularly cumbersome for ASIC lines where shipping dates and product types have constantly changing priorities.

JIT-CAM systems offer the advantage of knowing where all the lots are in the production cycle. The computer can stage the work to minimize disruption and keep a steady flow of product. Furthermore, JIT-CAM systems teamed with automatic delivery systems can be more efficient when the process stations are grouped rather than strung out in a line. This concept, called *work cells*, is more efficient, especially when a cell can run a number of different products. When a machine is inoperable in a linear layout, the line comes to a halt, and inventory builds up in front of the machine that is down.

Quality Control and Certification—ISO 9000

SPC and other product, process, or personnel quality programs fall under the general category of quality control. In the semiconductor industry, there are two broad quality segments: quality control (QC) and quality assurance (QA). QC generally refers to all of the techniques used to monitor and control the *process* and wafer quality while they are in the process. QA refers to the efforts to monitor and assure that the customer is receiving product that meets the required specifications.

In 1987, a system of quality requirements for organizations was introduced: the ISO 9000 series. Developed by the International Organization for Standardization (ISO), the requirements are more like guidelines, sometimes called an “umbrella standard,”¹⁹ rather than a strict cookbook set of standards.

The ISO 9000 guidelines/standards guide the development and critique of a complete quality system for product, process, and management. Individual companies decide on how best to comply and implement programs. With certification, customers and vendors know that there is a comprehensive quality program backing the product. The European Community (EC) requires that all companies doing business in the EC have ISO 9000 certification. In the United States, the American National Standards Institute (ANSI) and American Society for Quality Control (ASQC) have developed equivalent standards. Assisting semiconductor suppliers is a program sponsored by the Semiconductor Equipment and Materials International (SEMI).²⁰

Line Organization

Most fabrication areas are organized around the product-line concept. In this concept, fabrication areas are built to accommodate products with similar processing needs. Thus, there are bipolar lines and CMOS lines, and so forth. This arrangement makes for more efficient processing, since most of the machines are in use most of the time, and the staff can gain experience in processing a few products.

The staffs of these lines are also fairly self-contained. The primary responsibility falls to the fabrication or product manager ([Fig. 15.18](#)). Reporting to this individual are an engineering supervisor, a production manager (or general supervisor), a design department, and the equipment-maintenance group. The production manager is responsible for producing the finished wafers to specification, to cost, and to schedule. The engineering group is responsible for the developing of high-yield processes, documentation of the processes, and the daily sustaining of the line process. Both the production and engineering staffs are divided into groups focusing on a particular part of the process. This organization has the virtue of high focus on the fabrication area’s primary goal of producing chips at a profitable level.

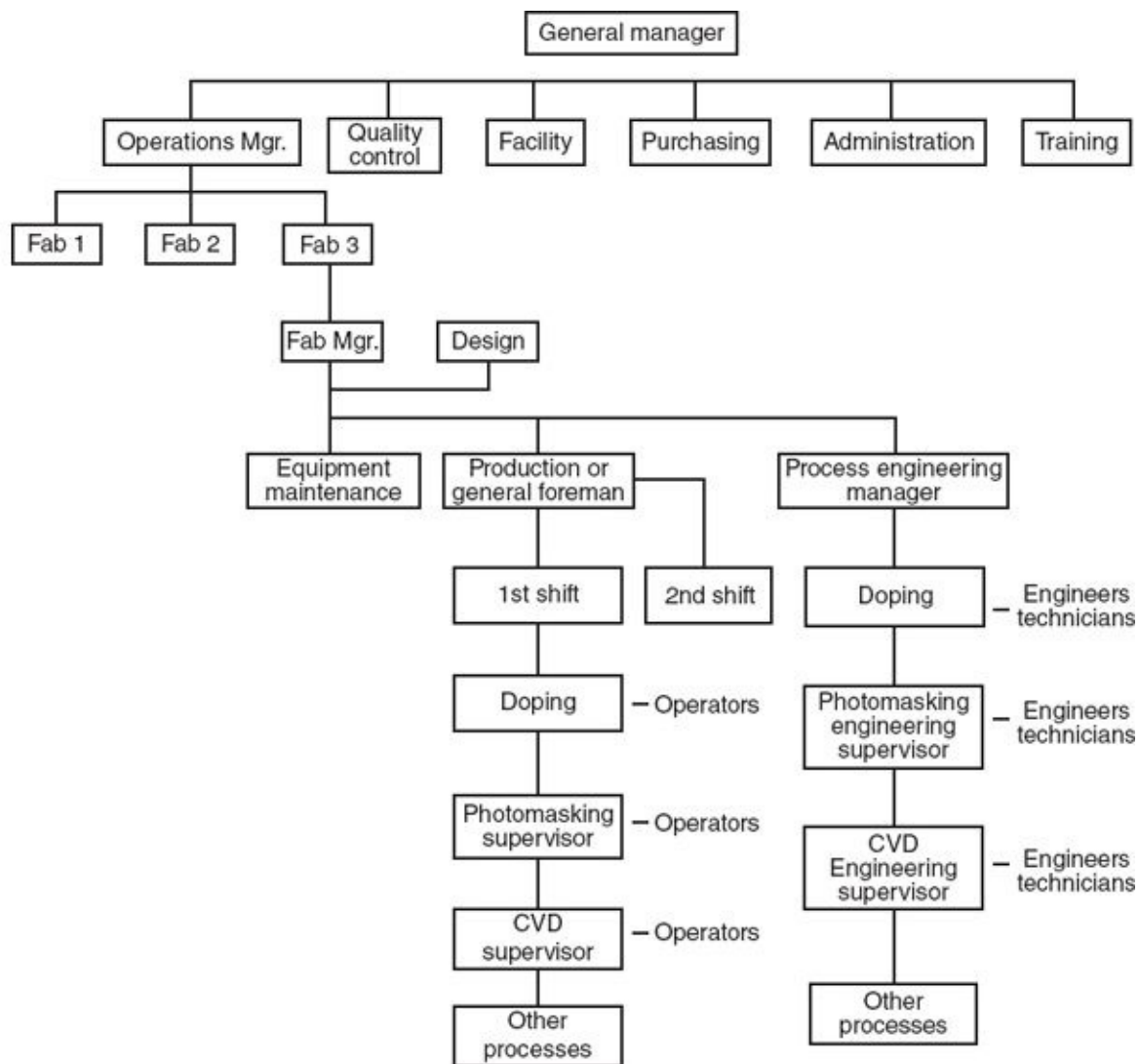


FIGURE 15.18 Typical semiconductor product line organization.

As the processes become more automated and arranged in process cells, small group organizational teams and responsibilities are emerging. A cell is attended by the operator(s), the equipment technicians, and the process engineer(s). These small groups make floor-level decisions with the information provided by the CIM system. However, few companies have formalized this arrangement with an organization structure, and the teams tend to exist as cross-department cooperatives.

Review Topics

Upon completion of this chapter, you should be able to:

1. List the major cost factors that influence fabrication costs.
2. Describe the intent and factors of cost of ownership (CoO) models.
3. List the advantages of statistical process control.

4. Identify the parts and use of a control chart.
5. List and discuss the different levels of automation.
6. List the factors that enter into an evaluation of a particular piece of equipment.
7. Define the terms “CIM” and “CAM” and their use in a manufacturing setting.

References

1. Semi/Sematech, *Semi Reports*, 2012, www.semi.org, May 2013.
2. Clark, P., “GlobalFoundries Hints at \$10 Billion Fab Location,” *EE Times*, Jan. 11, 2013.
3. Harper, J. G. and Bailey, L. G., “Flexible Material Handling Automation in Wafer Fabrication,” *Solid State Technology*, Jul. 1984:94.
4. Lam, D., “Minifabs Lower Barriers to 300 mm,” *Solid State Technology*, Jan. 1999:72.
5. Sonderman, T., *Reaping the Benefits of the 450-mm Transition*, Semicon West, San Francisco, CA:2011.
6. Arden, W., Brillouët, M., Coge, P., et al., *More Than Moore White Paper*, ITRS White Paper, Nov. 8, 2011.
7. Foster, L. and Pollai, D., *300-mm Wafer Fab Logistics and Automated Material Handling Systems, Handbook of Semiconductor Manufacturing Technology*, 2007, CRC Press, New York, NY:33–17.
8. Burggraaf, P., “Applying Cost Modeling to Stepper Lithography,” *Semiconductor International*, Cahners Publishing, Feb. 1994:40.
9. Shinoda, S., “Total Automation in Wafer Fabrication,” *Semiconductor International*, Sep. 1986:87.
10. Singer, P., “The Thinking behind Today’s Cluster Tools,” *Semiconductor International*, Aug. 1993:46.
11. Foster, L. and Pollai, D., *300-mm Wafer Fab Logistics and Automated Material Handling Systems, Handbook of Semiconductor Manufacturing Technology*, 2007, CRC Press, New York, NY:33–17.
12. Moslehi, M., “Single-Wafer Processing Tools for Agile Semiconductor Production,” *Solid State Technology*, PennWell Publications, Jan. 1994:35.
13. Sonderman, T., *Reaping the Benefits of the 450-mm Transition*, Semicon West, CA:2011.
14. Kerby, R. and Novak, L., “ESH: A Green Fab begins with You,” *Solid State Technology*, Jan. 1998:82.

[15.](#) Campbell, D. M. and Ardehale, Z., "Process Control for Semiconductor Manufacturing," *Semiconductor International*, Jun. 1984:127.

[16.](#) Levinson, W., "Statistical Process Control in Microelectronics Manufacturing," *Semiconductor International*, Cahners Publishing, Nov. 1994:95.

[17.](#) Levy, K., "Productivity and Process Feedback," *Solid State Technology*, Jul. 1984:177.

[18.](#) Ibid.

[19.](#) Hnatek, E., "ISO 9000 in the Semiconductor Industry," *Semiconductor International*, Jul. 1993:88.

[20.](#) Dunn, P., "The Unexpected Benefits of ISO 9000," *Solid State Technology*, Mar. 1994:55.

Introduction to Devices and Integrated Circuit Formation

Introduction

Integrated circuits are composed of individual conductors, fuses, resistors, capacitors, diodes, and transistors. The operation and formation of the basics of each is explored in this chapter as is the formation of the major integrated circuits from the components. An introduction to circuits is in [Chap. 17](#).

Semiconductor-Device Formation

The previous chapters have focused on the individual processes used to make semiconductor devices (also referred to as *components* or *circuit components*) and integrated circuits. It is assumed that the reader has already read about (or is familiar with) the processes and has a good understanding of the basic structure and electrical performance of the individual components as explained in [Chap. 14](#). There are literally thousands of different semiconductor-device structures. They have been developed to achieve specific performances, either as discrete components or in integrated circuits. However, there are basic structures required for each of the major device and circuit types. In this chapter, these basic structures are examined. Mastering them is essential to understanding the many variations and innovative structures that abound in the semiconductor world. The circuit components are: •

Resistors

- Capacitors
- Diodes
- Transistors
- Fuses
- Conductors

Resistors

Resistors have the effect of limiting current flow. This is accomplished by the use of dielectric materials or high-resistivity portions of a semiconductor wafer surface. In semiconductor technology, resistors are formed from isolated sections of the wafer surface, doped regions, and deposited thin films.

The value of a resistor (in ohms) is a function of the resistivity of the resistor and its dimensions ([Fig. 16.1](#)). The relationship is:

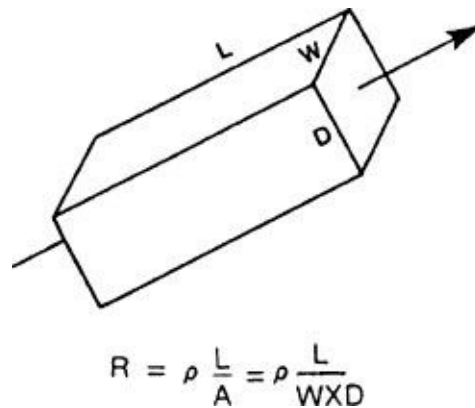


FIGURE 16.1 Relationship of resistance to resistivity and dimensions.

$$R = \rho LA$$

where ρ = resistivity

L = length of resistive region

A = cross-sectional area of the resistive region

The area (A) becomes $W \times D$, where W = width of the resistor and D = depth of the resistive region. For doped resistors, the length and width are the surface-pattern openings and the depth is the junction depth.

It should be obvious that every doped region is also a resistor, and the basic resistor formula governs electrical flow. A conductor is simply a resistor with a low resistance. The conceptual importance of Ohm's law is that the electrical resistance of any region in the device or circuit is altered by any change in dimensions or change in the doping level (resistivity).

Doped Resistors

Most of the resistors in integrated circuits are formed by a sequence of an oxidation, masking, and doping operation ([Fig. 16.2](#)). A pattern is opened in the surface oxide. Typical resistor shapes are dumbbells ([Fig. 16.3](#)) with the square ends serving as contact regions and the long skinny region in between serving as the resistor function. The resistance of this region is calculated from the sheet resistance of the region and the number of squares ($\square$ s) contained in the region. The number of squares is calculated by dividing the length by the width.

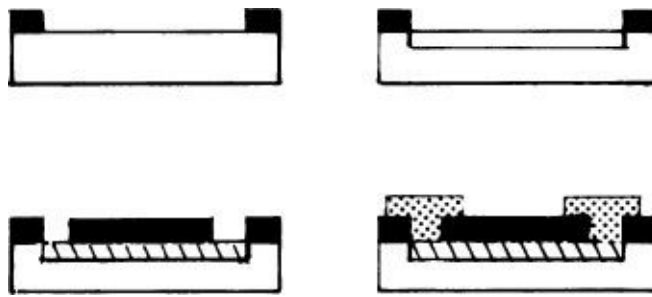


FIGURE 16.2 Diffused resistor formation.

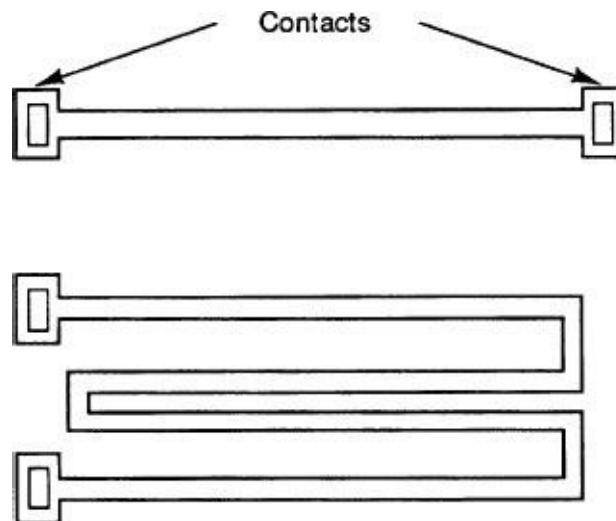


FIGURE 16.3 Resistor shapes.

After doping and a subsequent reoxidation, contact holes are etched in the square ends to contact the resistor into the circuit. A resistor is a two-contact, no-junction device. The term *no-junction* means that the current flows between the contacts without crossing an N-P or P-N junction. But the junction serves to confine the current flow in the resistive region.

Resistors doped by ion implantation have more controlled values than those in

diffused regions. Doped resistors can be formed during any of the doping steps performed during the fabrication process. A bipolar base mask will have the base pattern and a set of resistor patterns. In MOS circuits, resistors are formed along with the source-or drain-doping step. The resistor has the same doping parameters (sheet resistance, depth, and dopant quantity) as the transistor part. In these schemes the contacts to the resistors are formed after all of the other chip components (layers) are fabricated on the wafer.

EPI Resistors

A resistor can be formed by isolating a section of an epitaxial region ([Fig. 16.4](#)). After surface oxidation and contact-hole masking, what is left is a three-dimensional region functioning as a resistor.

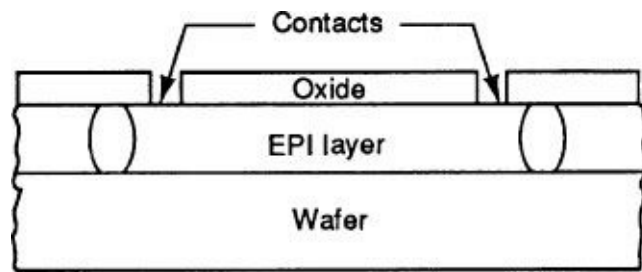


FIGURE 16.4 Epitaxial layer resistor.

Pinch Resistors

Ohm's law shows that the cross-sectional area of the resistor is a factor in its value ([Fig. 16.5](#)). One way to reduce the cross-sectional area (and increase the resistance) is to dope the resistor region and then do another doping of the opposite conductivity type. This occurs in bipolar processing when a resistor region is formed during the P-type base doping with a "pinched" cross-section formed from a subsequent N-type region formed along with emitter doping.

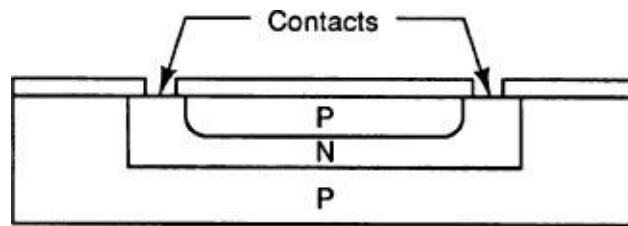


FIGURE 16.5 N-type resistor "pinched" with P-type doped region.

Thin-Film Resistors

Doped resistors do not always have the resistance control needed in some circuits and are poor performers in radiation environments. Radiation, such as found in space, generates unwanted holes and electrons that allow the current to leak across the confining junction. Resistors formed from deposited thin films of metal do not have this radiation problem.

The film resistors ([Fig. 16.6](#)) are formed from depositions of the film materials and patterning into the correct shape. After resistor formation of the wafer surface, it is "wired" into the circuit by contact between the resistor ends and the leads of the conducting metal. Nichrome, titanium, and tungsten are typical resistor metals.

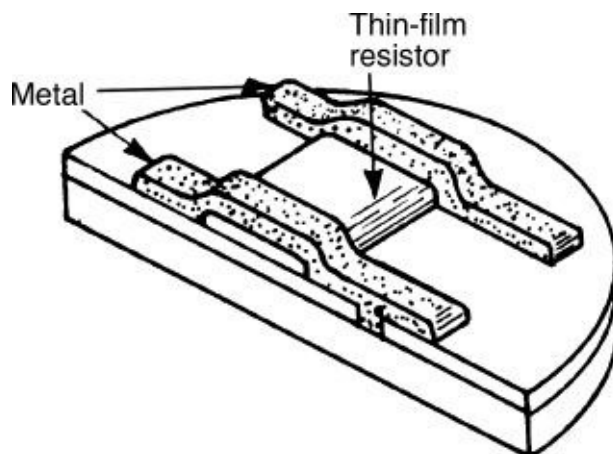


FIGURE 16.6 Formation of a thin-film resistor.

Capacitors

Oxide-Silicon Capacitors

Capacitors are formed from three layers. A dielectric layer sandwiched between two electrodes. Silicon planar technology is based on a silicon wafer with a grown silicon dioxide layer on top or on a silicon dioxide layer grown and an epitaxial layer. If a conducting metal lead lies on top of the oxide, a simple capacitor is formed ([Fig. 16.7](#)). Recall that a capacitor is formed with a dielectric layer sandwiched between two electrodes. In fact, this structure is a metal-oxide-metal (MOS) capacitor structure. However, the oxide thickness has to be thin enough (about $1500 \text{ \AA}$)¹ for the structure to act as a capacitor. The top electrode is also called a *cell plate*. A bottom electrode is also called the *storage node*.

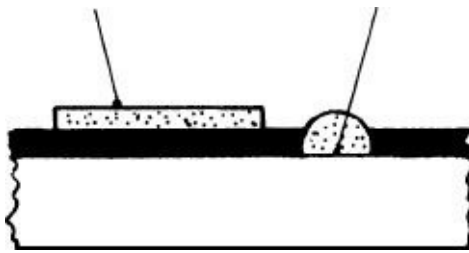


FIGURE 16.7 Monolithic capacitor.

A capacitor is a device that stores charge. A battery is a capacitor. When a voltage is applied to the metal, a charge builds up in the wafer layer under the oxide ([Fig. 16.7](#)). The amount of charge is a function of the oxide thickness, the dielectric constant of the oxide, and the area, as defined in the metal top plate. Capacitors of this structure are called parallel-plate, monolithic, or MOS (after the metal-oxide semiconductor materials of the sandwich).

In dense circuits, an oxide-nitride-oxide (ONO) dielectric sandwich is used. The combination film has a lower dielectric constant, allowing a capacitor area smaller than a conventional silicon dioxide capacitor. Capacitor function acts two ways in integrated circuits (ICs). In some circuits, capacitors are formed specifically to store charge. However, capacitor structures are formed whenever metal lines lie over a dielectric on top of a layer of silicon (or other semiconductor material). However, we do not want the capacitor to store charge that may interfere with the circuit operation. In this function, the dielectric layer needs to be thick enough to prevent capacitor charge storage. Also the use of high dielectric materials (high-k) can also prevent the top-side conductor from building up a charge in the base silicon. See [Chap. 12](#) for reference.

Junction Capacitors

A capacitor is formed at every junction in a device. When a voltage is applied across any junction, carriers on each side of it move away from the junction, leaving a depleted region (Fig. 16.8). This depleted region acts as a capacitor in the device or circuit.

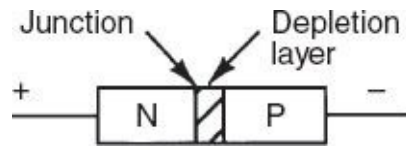


FIGURE 16.8 Depletion layer junction capacitor.

The value of this junction capacitance must be taken into account when the circuit is designed. Some circuits actually use junction capacitors as part of the circuit design. In some circuits, the natural junction capacitance has the effect of slowing up the circuit operation. This is due to the time required to “fill up,” or charge, the depleted region before current flows. A finite time is also required for the various junction capacitors to discharge. Both of these times affect circuit switching and operational speeds.

Trench Capacitors

Preservation of wafer surface area is always a design criterion. One of the problems with oxide-metal capacitors is their relatively large area. Trench (or buried) capacitors solve the problem by creating a capacitor in a trench etched vertically into the wafer surface ([Fig. 16.9](#)). The trenches are etched either isotropically with wet techniques or anisotropically with dry etch techniques. The trench sidewalls are oxidized (the dielectric material) and the center of the trench is filled with deposited polysilicon. The final structure is “wired” from the surface, with the silicon and polysilicon serving as the two electrodes with the silicon dioxide dielectric between them. Other dielectric materials may be used in place of the silicon dioxide to increase performance.

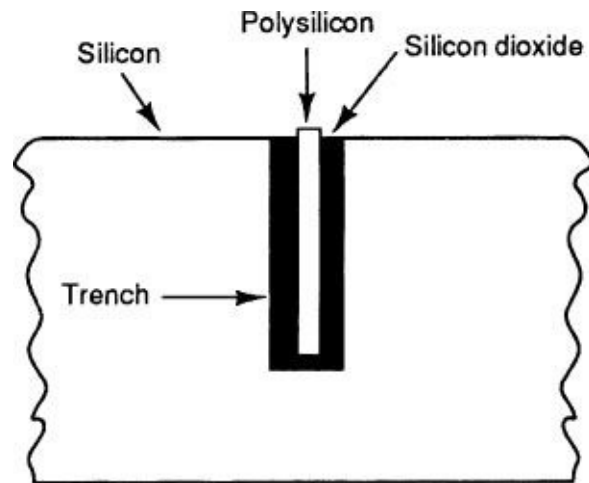


FIGURE 16.9 Trench capacitor.

Stacked Capacitors

Another alternative for conserving surface real estate is to build *stacked capacitors* on the wafer surface. This effort has been driven by the need for small, high-dielectric capacitors for dynamic random access memory (DRAM) circuits. The storage part of a DRAM cell is a capacitor. In DRAM cells, the bottom electrode is usually polysilicon or hemispherical grain polysilicon (HSG).²

Capacitor dielectric materials include tantalum pentoxide (Ta_2O_5) and barium strontium titanate (BaSrTiO_3 or BST). The latter is a *ferroelectric* material.³ Ferroelectric refers to iron-containing materials and, in the world of electronics, represents dielectrics with improved speeds over conventional silicon-compatible materials. Another ferroelectric material is $\text{PbZ}_{1-x}\text{T}_x\text{O}_3$ or PZT.

The top electrode materials may be TiN, WN, Pt, polysilicon, or one of the other semiconductor conducting materials.

Diodes

Doped Diodes

A diode is a two-region (two-contact) device separated by a junction. A diode either allows current to pass easily or acts as a current block. Which function it performs is determined by the voltage polarity, called *biasing* ([Fig. 16.10](#)). When the current voltage is the same as in the diode region, the diode is in forward-bias, and the current flows easily. When the polarities are reversed, the diode is reverse-biased, and the current is blocked. A reverse-biased diode can be forced into a conducting state by raising the current voltage until the junction goes into *breakdown*. This condition is temporary; when the voltage is reduced, the diode once again becomes a blocking device (see [Chap. 14](#)). Diodes are used to steer the current around a circuit. By proper choice of the circuit current polarities and the correct diode polarities, the current is allowed to pass into some branches of the circuit and is blocked out of others. A planar diode is formed from a doped region and two contacts on either side of the junction where it intersects the surface ([Fig. 16.11](#)). Diodes are usually formed along with transistor-doping steps. Thus, in bipolar circuits, there are base-collector diodes and emitter-base diodes. In MOS circuits, most of the diodes are formed with the source-drain doping step.

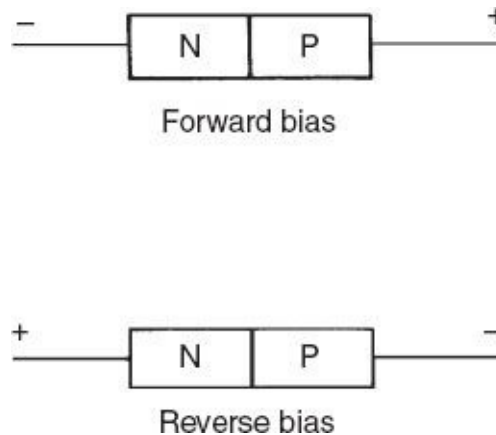


FIGURE 16.10 Forward and reverse biasing.

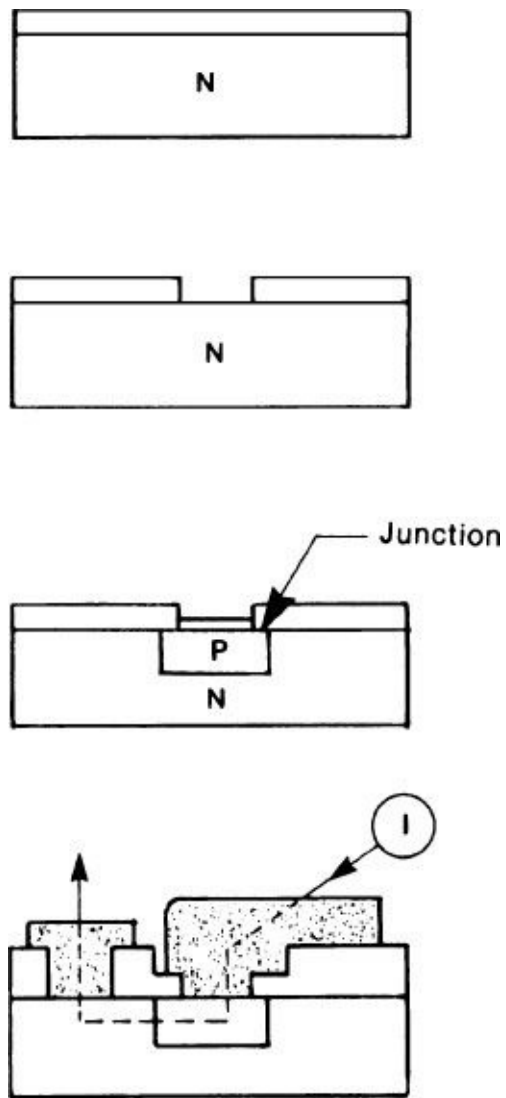


FIGURE 16.11 Formation of a P/N planar diode.

Schottky Barrier Diodes

In 1938 (10 years before the invention of the transistor), W. Schottky⁴ discovered that whenever a metal is in contact with a lightly doped semiconductor, a diode is formed ([Fig. 16.12](#)). This diode has a faster forward time (it responds faster) and operates with a lower voltage than a doped silicon junction diode. Metal contacts to highly doped regions (greater than 5×10^{17} atoms/cm³) are regular ohmic contacts. This is the situation for the majority of contacts in a silicon circuit. This Schottky diode effect is taken advantage of in some NPN bipolar transistors. The structure and effect are explained in the section on bipolar transistors.

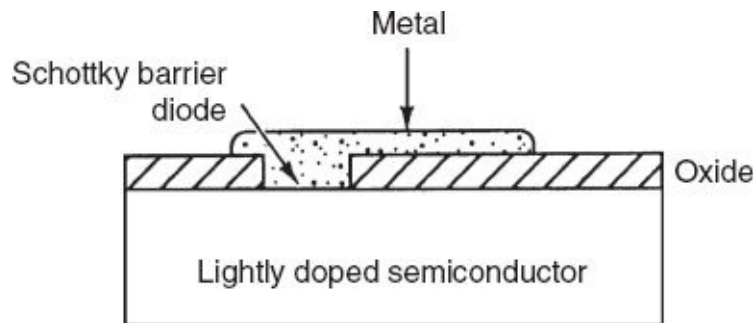


FIGURE 16.12 Schottky barrier diode.

Transistors

Transistor Operational Analogy

A transistor is a three-contact, three-part, two-junction device that performs as a switch or an amplifier. An often used analogy to explain the role of the parts and the operation of a transistor is the water-flow system in [Fig. 16.13](#). The flowing water represents current flow. In this system, one part is the source of the water (the tank), the valve controls the flow, and the bucket collects the water. The system can be operated as a switch simply by turning the valve on and off. It can even be imagined in an amplifier role. Consider the valve as a high-mechanical advantage miniature water wheel activated by a small external stream of water. A small trickle onto the valve wheel could open the valve to allow a large flow through the system. If the whole system was enclosed so that an observer saw only the trickle going in and a large flow coming out, such an observer might conclude that the system was *amplifying* the water trickle. Transistors are formed to provide the same functions as described below.

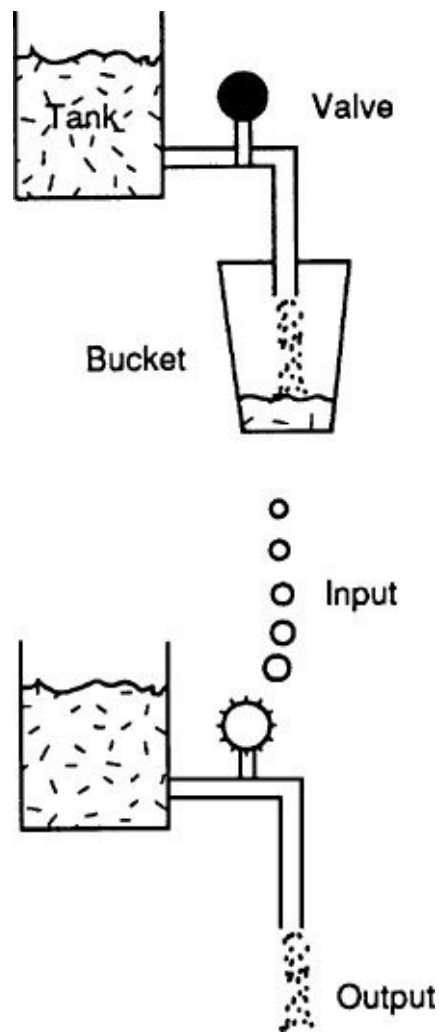


FIGURE 16.13 Water analogy of transistor operation.

William Shockley noted the function of a field effect transistor (FET) in 1948. Yet the first several decades after the discovery of the transistor at Bell Labs, the bipolar transistor became the dominate transistor structure. The advent of large-scale integration in the 1960s favored the FET in the metal-oxide silicon (MOS) structure (MOSFET). Transistors operate in both the NPN and PNP configurations. Engineers were able to fabricate each (NPN and PNP) into circuits known as complementary metal-oxide silicon (CMOS). These devices have dominated the industry because of their function in digital circuits and the ability to maintain functionality as they have been scaled to ever-smaller dimensions. It is scaling that allowed component reduction and higher circuit device count per Moore's law, in planar technology. Intel moved into the third dimension with the development of the tri-gate transistor also named the FinFET. These structural designs are described in the following sections.

Bipolar Transistors

The same basic parts and functions are present in solid-state transistors. A bipolar transistor is shown as both a simple block arrangement and in a double-doped planar form in [Fig. 16.14](#). The current flows from the emitter region (tank) through the base (valve) into the collector (bucket). When there is no current to the base, the transistor is turned off. When it is on, the current flows. It only takes a small current to turn the base on enough to allow current flow through the whole transistor. The size of the base current regulates the larger amount of current flowing through the transistor (called the *collector current*). This is an amplification of the base current to the collector current. The base current in effect changes the resistance of the base region. In fact, the term *transistor* comes from an early term for a bipolar transistor: *transfer resistor*. During operation, both positive and negative currents flow in the base, hence the term *bipolar*—literally, two polarities.

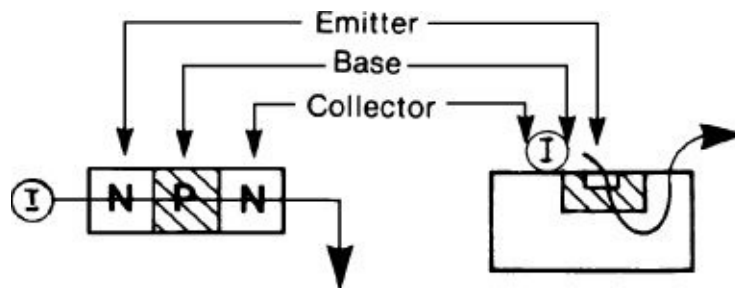


FIGURE 16.14 Bipolar transistor operation.

The amplification property, called *gain* or *beta*, is numerically the result of dividing the collector current by the base current (see [Chap. 14](#)). For efficiency, the emitter region has a higher doping concentration than the base, and the base doping is higher than the collector. A typical doping concentration versus distance plot is shown in [Fig. 16.15](#).

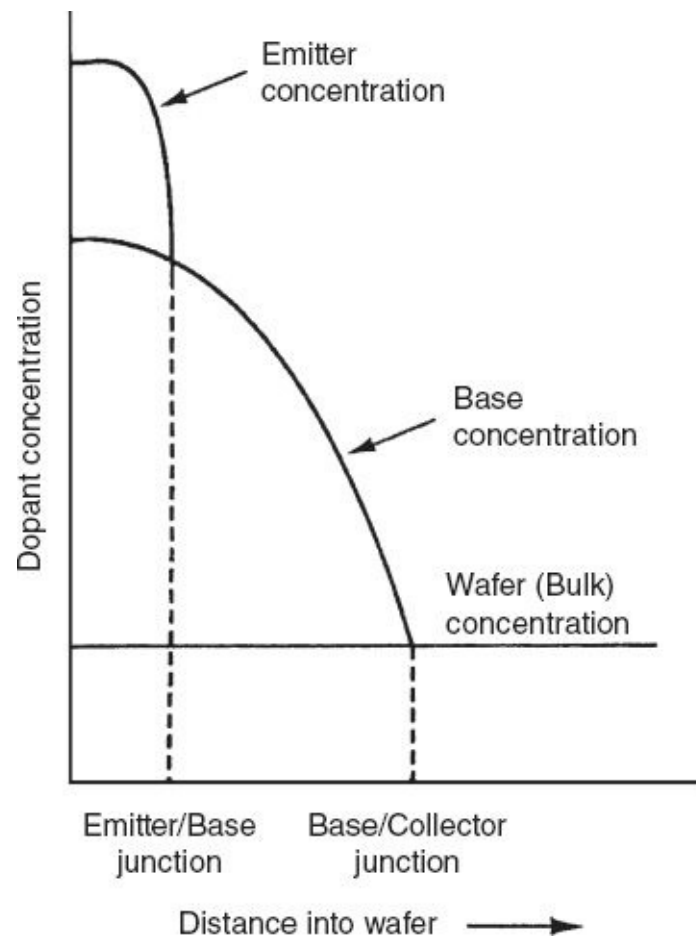


FIGURE 16.15 Dopant concentration profile of bipolar transistor.

Most bipolar circuits are designed with NPN transistors. NPN represents the respective conductivity types of the emitter, base, and collector. Some applications require PNP transistors, with many of them being formed laterally ([Fig. 16.16](#)). NPN transistors are more efficient because of the ease (higher mobility) of electron movement in the N-type regions.

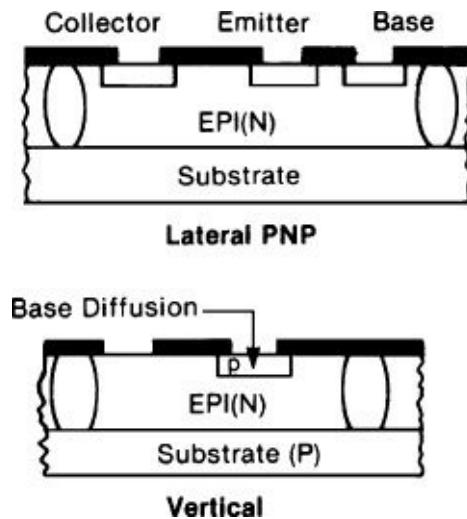


FIGURE 16.16 Lateral and vertical PNP transistors.

Bipolar transistors feature fast switching speeds. The speed is governed by a number of factors, of which the most important is the base width. Applying commonsense, it stands to reason that the shorter the distance an electron or hole has to travel, the less time it will take.

Bipolar transistors can switch on and off in as little as a billionth of a second. To achieve this speed, the transistor is maintained at an “on” state. This means that the base always has power applied, which is a downside of bipolar-based circuits. Another penalty for this necessary condition is a buildup of heat in the transistor. This heat affects circuit operation and is the reason for the cooling fans and air-conditioned rooms of earlier bipolar-based computers.

Schottky Barrier Bipolar Transistors

The Schottky barrier diode principle mentioned in the previous “Diodes” section is put to use in some bipolar transistors ([Fig. 16.17](#)). The construction requires that the base contact be extended into the collector region. When covered with metal, a Schottky diode is formed between the base and collector, which results in a faster-responding transistor. In a circuit, the time required to turn the transistor on and off (switching speed) is critical when millions of transistors are operating.

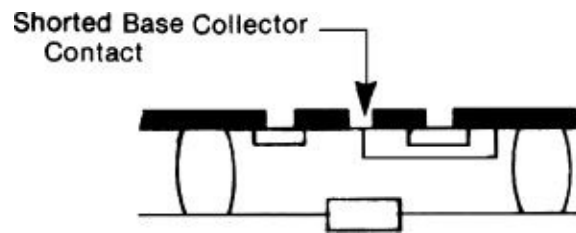


FIGURE 16.17 Schottky barrier bipolar transistor.

Field-Effect Transistors

Metal-Gate MOSFET

Switching and amplification are also achieved in an FET transistor with the MOS structure as the most popular design. A MOSFET transistor, like a bipolar ([Fig. 16.18](#)), has three regions, three contacts, and two junctions, but in a different structure. There is a similar analogy to the water system described previously. Current travels from the source region (tank), through the dielectric gate material (valve), and into the drain (bucket) before exiting the device.

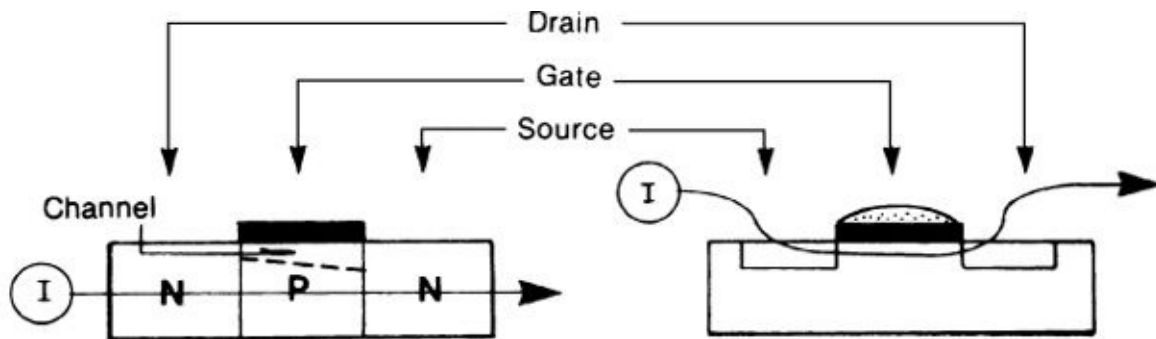


FIGURE 16.18 MOS transistor operation.

A MOSFET gate controls current flow by a different mechanism from that of a bipolar base. The MOS structure shown in [Fig. 16.19](#) is a simple metal-gate type

capacitor and operates the same as a capacitor. When a voltage (the gate voltage) is applied to the gate through the gate metal, a *field effect* takes place in the surface of the semiconductor. The effect is either a buildup of charge or a depletion of charges in the wafer surface under the top plate. Which event occurs depends on the doping conductivity type in the wafer under the gate and the polarity of the gate voltage.

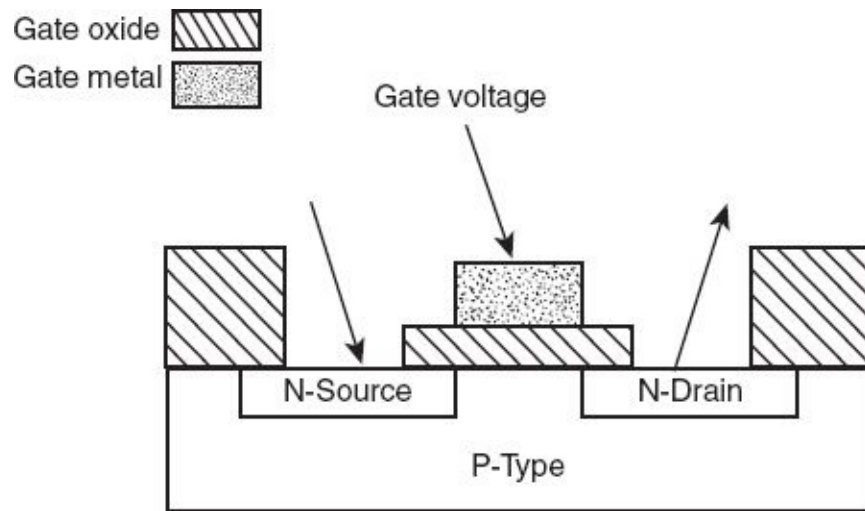


FIGURE 16.19 Metal gate MOS transistor.

The buildup or depletion of charge creates a channel under the gate that connects the source and drain. The surface of the semiconductor is said to be *inverted*. The source is biased with a voltage, and the drain is grounded relative to the source. In this condition, a current starts to flow as the inverted surface creates an electrical connecting channel. The source and drain are essentially shorted together. Applying more voltage to the gate increases the size of the channel, allowing more current to flow through the transistor (see [Chap. 14](#)). By controlling the gate voltage, a MOS transistor can be used as a switch (on/off) or as an amplifier. However, MOS transistors are voltage amplifiers, unlike the current amplification of bipolar transistors.

If the source and drain are N-type formed in a P-type wafer, the channel must be of N-type for conduction to occur. This type of MOS transistor is called N-channel. MOS transistors with P-type sources and drains are P-channel. Most high-performance MOS circuits are built around N-channel transistors due to the higher mobility of electrons in N-type silicon. The mobility makes N-channel transistors faster, and they consume less power than P-channel circuits. They often are referred to as NMOS transistors. [Figure 16.20](#) shows the major steps in the formation of an N-channel metal-gate MOS transistor.

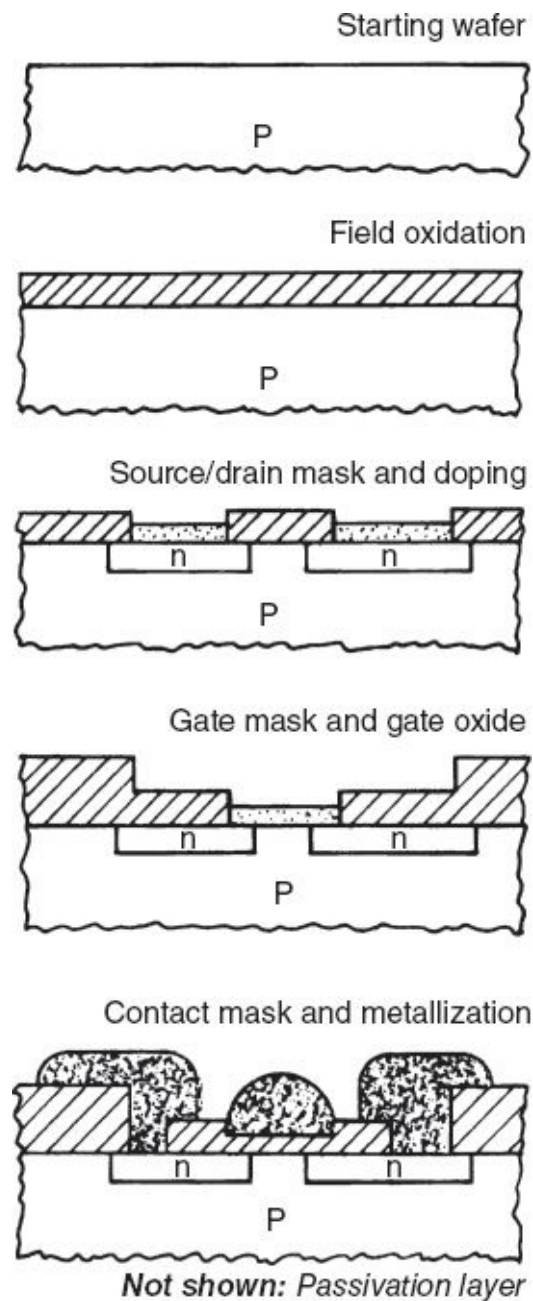


FIGURE 16.20 N-channel metal gate MOS process.

Silicon-Gate MOS

A certain amount of voltage must be applied to the gate metal before the channel forms. This voltage is called the *threshold voltage* or V_t . The value of the threshold voltage is an important and critical circuit parameter. A lower V_t means fewer power supplies and faster circuits.

A primary parameter that determines the threshold voltage is the *work function* between the gate material and the doping level in the semiconductor. The work function can be thought of as a kind of electrical compatibility. The lower the work function, the lower the threshold voltage, the lower the power required to run the

circuit, and so on.

Deposited doped polysilicon has a lower work function than aluminum as an MOS gate material and has become the standard gate electrode material for MOS transistors. The formation of the transistor is shown in [Fig. 16.21](#). The polysilicon is a heavily doped N-type to reduce its resistance. Thus doped, it serves as the gate electrode and as a circuit conduction line. A polysilicon gate can withstand subsequent high-temperature processing without degradation.

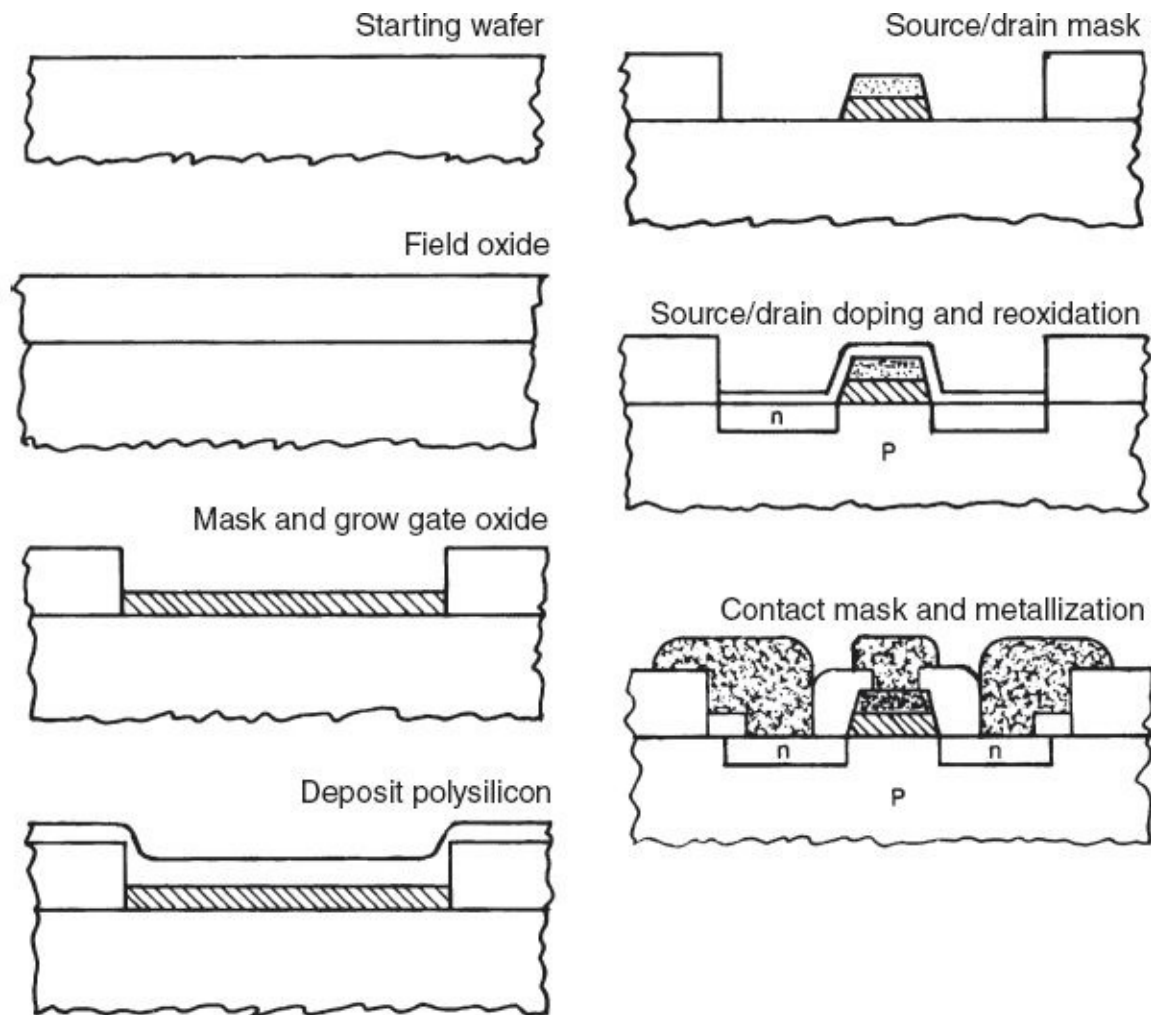


FIGURE 16.21 Silicon gate MOS process steps.

An additional benefit of the silicon-gate process is the self-aligned gate. In the metal-gate process sequence, a hole for the gate oxide must be patterned between the source and drain. To ensure that the source and drain are bridged by the gate, overlap for alignment tolerances must be allowed. This results in some overlap of the gate into the source and/or drain. The overlap becomes a region of unwanted capacitance. In the silicon-gate process, the gate is formed first and acts as a mask to locate the source and drain. Thus, whatever the gate placement, the source and drain

self-align to it.

As MOSFET technology has developed, the gate material has evolved. The trend has been to lower gate threshold voltage and thinner dielectric thickness and smaller gate areas. Unfortunately, thinner and smaller gates also increase gate leakage. Silicon oxynitride films (SiON) were an improvement on simple SiO₂ films. However, SiON reaches a limit on minimizing gate leakage below gate thickness of 1.0 for high-performance devices and 1.5 nm for low-power devices. Hence, the interest and high-dielectric (high-k) gate materials. The capacitance formula below

shows the relationship of the parameters: $C = \frac{kE_0 A}{t}$

where C = capacitance

k = dielectric constant of material

E_0 = permittivity of free space (free space has the highest “capacitance”)

A = area of capacitor

t = thickness of dielectric material

With gate thickness and area at the limits of existing technology the next parameter to change is the dielectric constant of the gate materials. That is why high- k materials are now part of the gate material advances.⁵

Other factors, besides the gate metal material, affecting the gate threshold voltage and operation are:

- Gate oxide thickness

- Gate material (dielectric constant)
- Source-drain separation (channel length)
- Gate doping level
- Sidewall capacitance of the doped source and drain regions

Gate thickness on the high- k materials is normalized to an equivalent silicon dioxide thickness (t_{ox}). However it is calculated, the actual dielectric layer is down to counting in number of atoms.

The thinner the gate oxide, the faster the device and the lower the threshold voltage. Oxide-nitride sandwiches (ONO), oxide tantalum pentoxide (OTa₂O₅), and amorphous silicon film stacks are used.

Channel lengths also affect the speed. Channel lengths have continually shrunk to the submicrometer range. In self-aligned structures, the channel length is established by the gate width. The gate-doping level influences the threshold voltage by modifying the work function difference between the gate metal and surface. Ion implantation, with its ability to dope through thin oxides, is often used to set the gate-doping level. Sidewall capacitance of the doped source and drain also serves to slow up the operation of the device, as enough charge must be built up to overcome the junction capacitance.

Polycide-Gate MOS

MOS development was initially impeded by the inability of the industry to grow non-contaminated and thin oxides in the 1960s. Contamination, especially the mobile ionic variety, interferes with the field effect, making very unreliable gates. In fact, clean gate oxides and the oxide-silicon interface are so well understood and effective that gate-oxide replacement is proceeding very cautiously.

Consequently, the quest for a higher-efficiency gate has led to the polycide structure ([Fig. 16.22](#)). The gate sandwich retains the thin oxide on the wafer surface topped by a layer of polysilicon. The polysilicon provides a low work function (and lower threshold) gate, and the reliable polysilicon-oxide interface is preserved. The new layer is a refractory metal silicide on top of the polysilicon. The silicide makes a low contact resistance with the polysilicon (as compared with aluminum) and reduces the overall sheet resistance of the polycide sandwich.

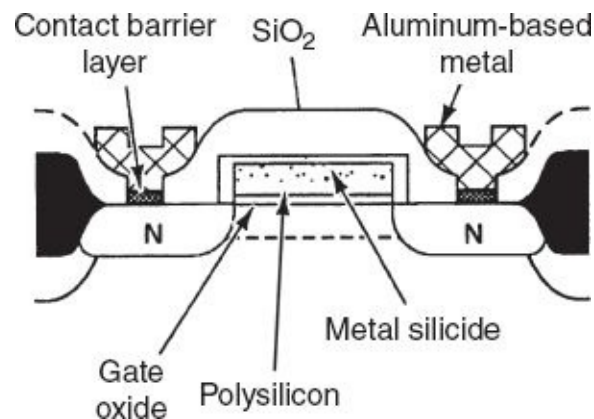


FIGURE 16.22 Polycide gate structure.

Salicide-Gate MOS

The self-aligned process that uses the polycide-gate structure is called a *salicide* gate. Its formation is illustrated in [Fig. 16.23](#). The process combines the best features of a polysilicon gate with self-alignment. The source and drain are lightly diffused around the polysilicon gate. Then a layer of silicon dioxide is deposited and anisotropically etched to form spacers on the side of the gate. These spacers act as ion-implantation masks for a subsequent heavier doping of the source and drain. The more lightly doped “finger” under the gate is called a lightly doped drain (LDD) extension.⁶ After ion implantation, the refractory metal is deposited, and the silicide is formed by a reaction with the underlying polysilicon layer by an alloy step. The final step is the removal of the unreacted refractory metal from the wafer surface.

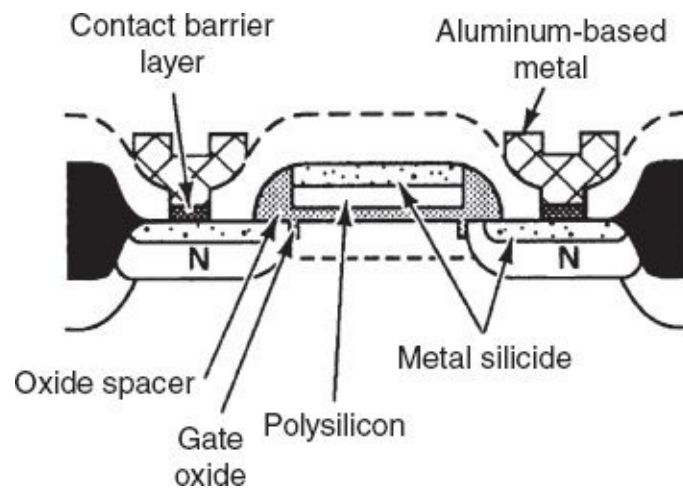


FIGURE 16.23 Salicide-gate structure.

The LDD process illustrated puts lightly doped fingers in both the source and drain areas. There are processes designed to create asymmetrical LDD structures that place the finger only in the drain region.²

Diffused MOS

Diffused MOS (DMOS) refers to a diffused MOS structure used in power MOSFETs. ([Fig. 16.24](#)). The channel length is established by two diffusions through the same opening. As the second diffusion is taking place, the first moves laterally to the sides. The second diffusion functions as the source, and the bulk semiconducting material of the wafer functions as the drain. The difference between the two diffusion widths is the channel length of the transistor.

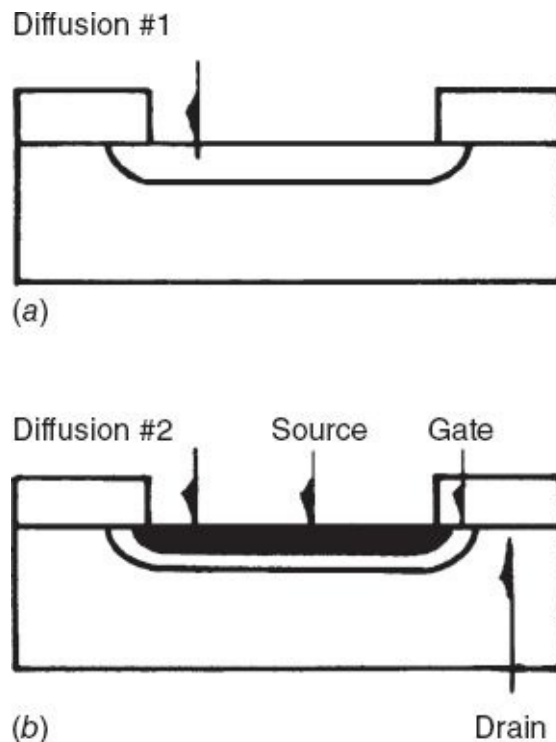


FIGURE 16.24 DMOS structure.

The double diffusion technique results in a vertical MOS transistor ([Fig. 16.24b](#)) with the P and N pockets in an epitaxial layer.

Memory MOS

Memory MOS (MMOS) is a structure that provides a more or less permanent storage of the charge in the gate region. The storage is provided by a thin layer of silicon nitride between the wafer and the gate oxide ([Fig. 16.25](#)). When the gate is charged to store data, the silicon nitride layer traps and retains it. This type of transistor is used in nonvolatile circuits where protection against memory loss is important (see [Chap. 17](#)).

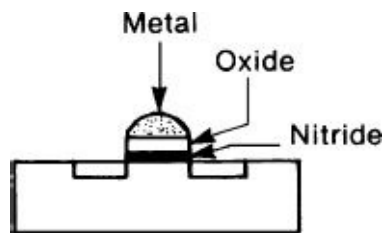


FIGURE 16.25 MMOS structure.

Junction Field-Effect Transistors

A junction field-effect transistor (JFET, [Fig. 16.26](#)) is similar in construction to a MOSFET but has a junction formed under the gate. During operation, the current flows *under* the diffused region from the source to the drain. As the gate voltage is increased, a region depleted of charge (the depletion region) spreads under the junction toward the N-type and P-type interface. The depleted region will not support current flow and has the effect of restricting the current flow as it increases in depth.

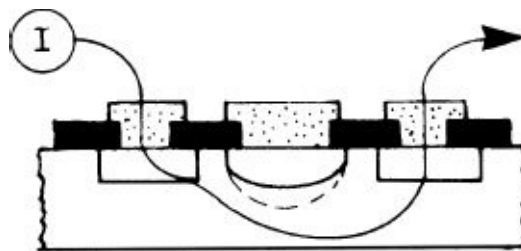


FIGURE 16.26 Junction field-effect transistor.

A JFET operates opposite from an MOS transistor. In the MOS version, increasing the gate voltage *increases* the current flow. In a JFET, increasing the gate voltage *decreases* the current flow. JFETs are a standard gallium-arsenide device. The transistor is formed in an N-type GaAs layer that is on top of a semi-insulating GaAs wafer ([Fig. 16.27](#)).

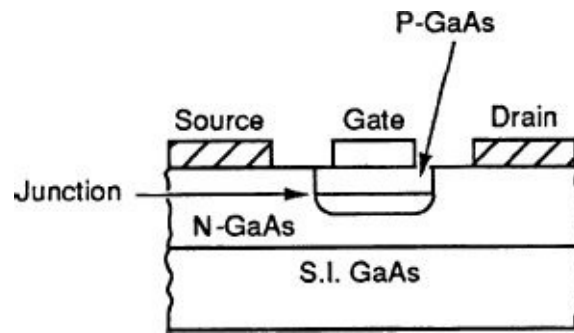


FIGURE 16.27 GaAs JFET.

Metal Semiconductor Field-Effect Transist

The metal semiconductor field-effect transistor (MESFET) is the basic GaAs transistor structure ([Fig. 16.28](#)). The MESFET operates in the same manner as the JFET, but the gate metal is deposited directly onto the N-type GaAs layer.

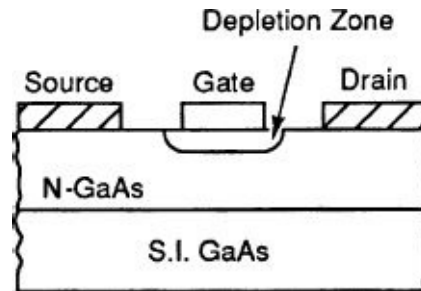


FIGURE 16.28 GaAs MESFET.

Alternatives to MOSFET Scaling Challenges

A standard approach to denser circuits is *scaling*, also called a *die shrink*. Scaling starts with a proven design and reduces the dimensions. However, scaling can introduce new problems, such as increased leakage. Many of the new materials are designed to address leakage and other scaling issues. Alternative transistor structures are also being explored. Ideas include double gates (DGs) and ultra-thin body (UTB) for MOSFET devices. Schematics are illustrated in [Fig. 16.29](#).

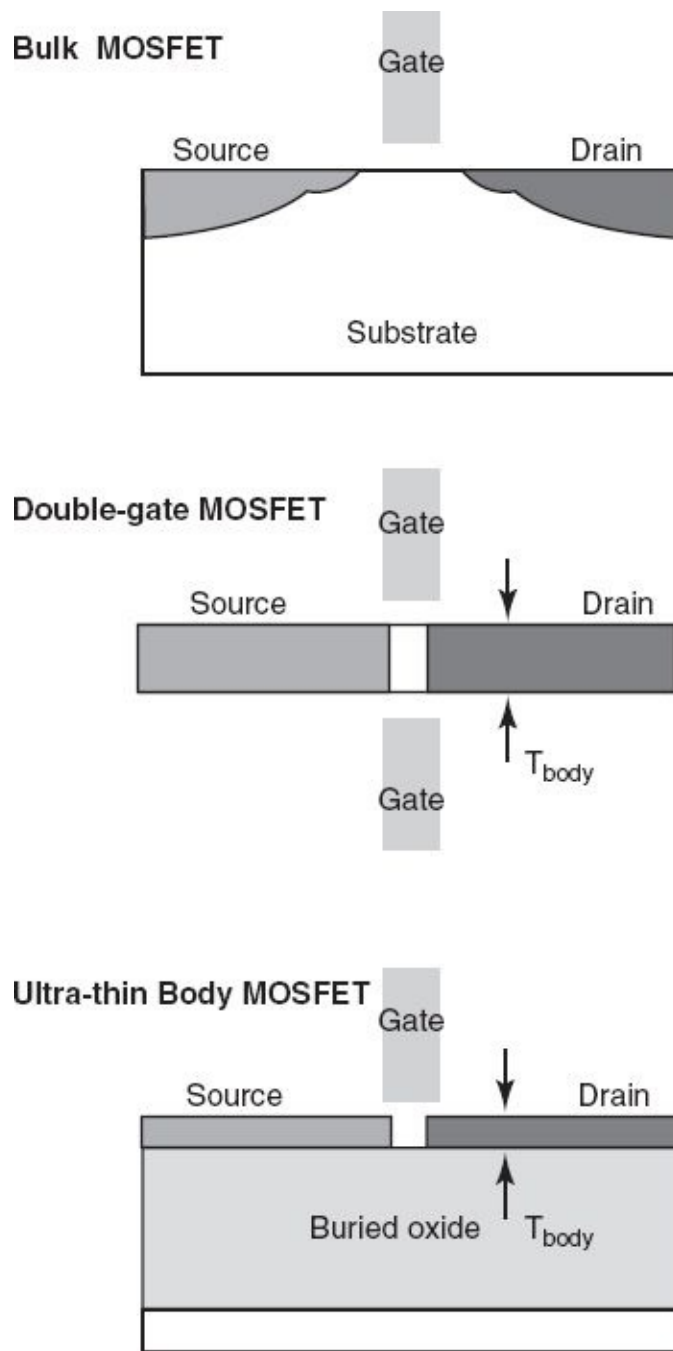


FIGURE 16.29 Alternate MOSFET devices. (Courtesy of Semiconductor International.)

Another evolution is a vertical transistor structure with a three-dimensional fin structure. Called a *FinFET* the gate is built vertically, and the source and drain build up on either side of the fin.

The advance is the raised fin covered with a high-k metal gate stack ([Fig. 16.30](#)).⁸ The electrical effect is an increased gate area from the two sides and top of the stack. FinFET transistors are built on either a bulk silicon substrate or on a silicon on insulator (SOI) substrate. In each case, the “fin” is created by etching down into the base material as illustrated in [Fig. 16.30](#). The gate stack is deposited and etched.

Intel pioneered the tri-gate structure in 2002 and it is becoming an industry standard. The FinFET is another solution to the limits of Moore's law by taking the structure into the vertical (third) dimension. And another innovation is multiple gates ([Fig. 16.31](#)) on the transistor structure.

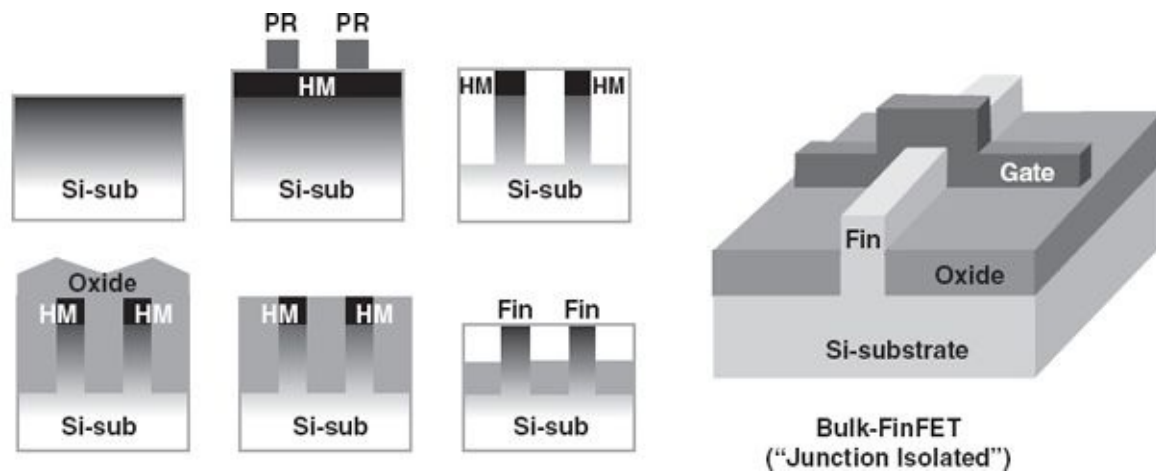


FIGURE 16.30 Process flow for bulk FinFET.

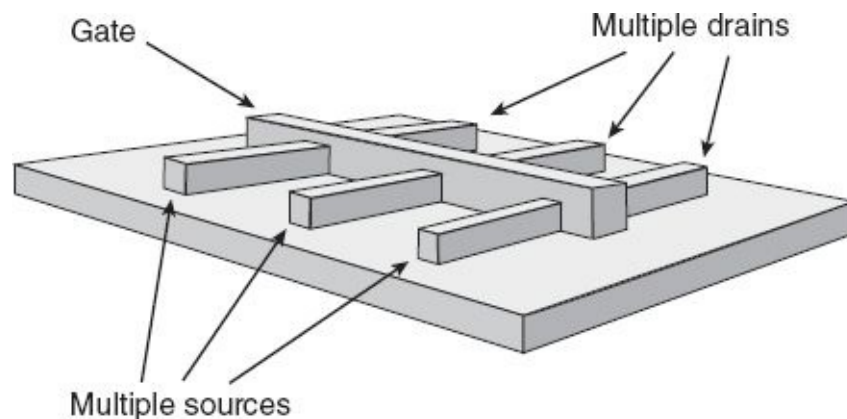


FIGURE 16.31 Multiple gate structure.

Conductors

The formation of surface wiring conductors were detailed in [Chap. 14](#). In very dense circuits, precious surface area is preserved by the use of subsurface *underpass conductors*. These are created from heavily doped regions formed under surface wiring leads.

Multiple metal layers have been described in other chapters as another way to achieve increased performance out of a given die area. The issues and challenges are depositing and patterning multiple metal layers. Key technologies are planarization, step coverage, and plug filling. Development of low-contact-resistance metal film and film sandwiches will be required in the 0.5- μm design rule era.

Integrated-Circuit Formation

Integrated circuits contain all the components described in the previous sections. The components are formed in specific sequences with the process flows designed around the transistor in the circuit. The process designer will attempt to have as many component parts as possible with each doping step.

Circuits are designated by the transistor type. A bipolar circuit means that the circuitry is based on bipolar transistors. MOS circuits are based on one of the MOS transistor structures. For the first 30 years of the semiconductor industry, the bipolar transistor and bipolar circuits were the structures of choice. Bipolar transistors had fast speeds (switching times), control of leakage currents, and a long history of process development. These qualities fit nicely into the logic, amplifying, and switching circuits that were the first offerings of the industry. These circuits handled the computational requirements of the growing computer industry. The internal memory functions of the early computers were handled by core memories. These memories were limited in capacity and slow. Much of the information needed was stored outside the computer on tape, disks, or punch cards. While bipolar memory circuits were available, they could not compete economically with core memories.

MOS transistor circuits held the promise of fast, economical, solid-state memory, but early metal-gate MOS circuits suffered from high leakage currents and poor parameter control. Even so, the built-in advantages of MOS transistors drove the development of MOS memory circuits. The advantages are smaller dimensions, which allow denser circuits, and relatively faster switching speeds. Yields tend to be higher on smaller-dimensioned circuits, since a given defect density will affect fewer transistors and components.

Perhaps the biggest density factor advantage of MOS components is the smaller area required for isolation of adjacent components. The various isolation schemes used are discussed in the following sections. Another advantage is low-power operation. First, MOS transistors sit in the circuit in the “off” mode, not soaking up power or generating heat like bipolar transistors, which must be on all the time to be in a “ready” state. Second, MOS transistors, being voltage-controlled devices, require a lower power to operate. CMOS circuits are an integrated circuit design that reduces power requirements even further.

An initial advantage of MOS circuits was fewer processing steps and smaller die sizes, which made for lower processing costs and higher yields. These advantages have disappeared as MOS circuits have evolved to VLSI/ULSI size with the additional steps required to fabricate CMOS circuits. In general, the faster switching speed of bipolar circuits has made them favored for logic circuits. MOS circuits, with their smaller component dimensions and lower power requirements, have been incorporated into memory circuits. By the 1980s, these traditional uses blurred, with CMOS technology being the preferred system for most circuit designs. These topics are addressed further in [Chap. 17](#).

Bipolar Circuit Formation

Junction Isolation

The bipolar transistor structure and basic performance have been illustrated. However, combining transistors with the other devices to form a circuit requires additional structures. They include an isolation scheme and a low-resistance collector contact. If two transistors or other devices are fabricated next to each other, they will not work, because they are electrically shorted together ([Fig. 16.32](#)). An early challenge of bipolar IC designers was finding a way to *isolate* the various circuit components. This need led to the epitaxial (epi) layer bipolar structure ([Fig. 16.33](#)).

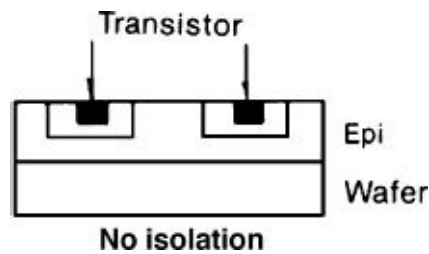


FIGURE 16.32 Adjacent bipolar transistors with common collectors.

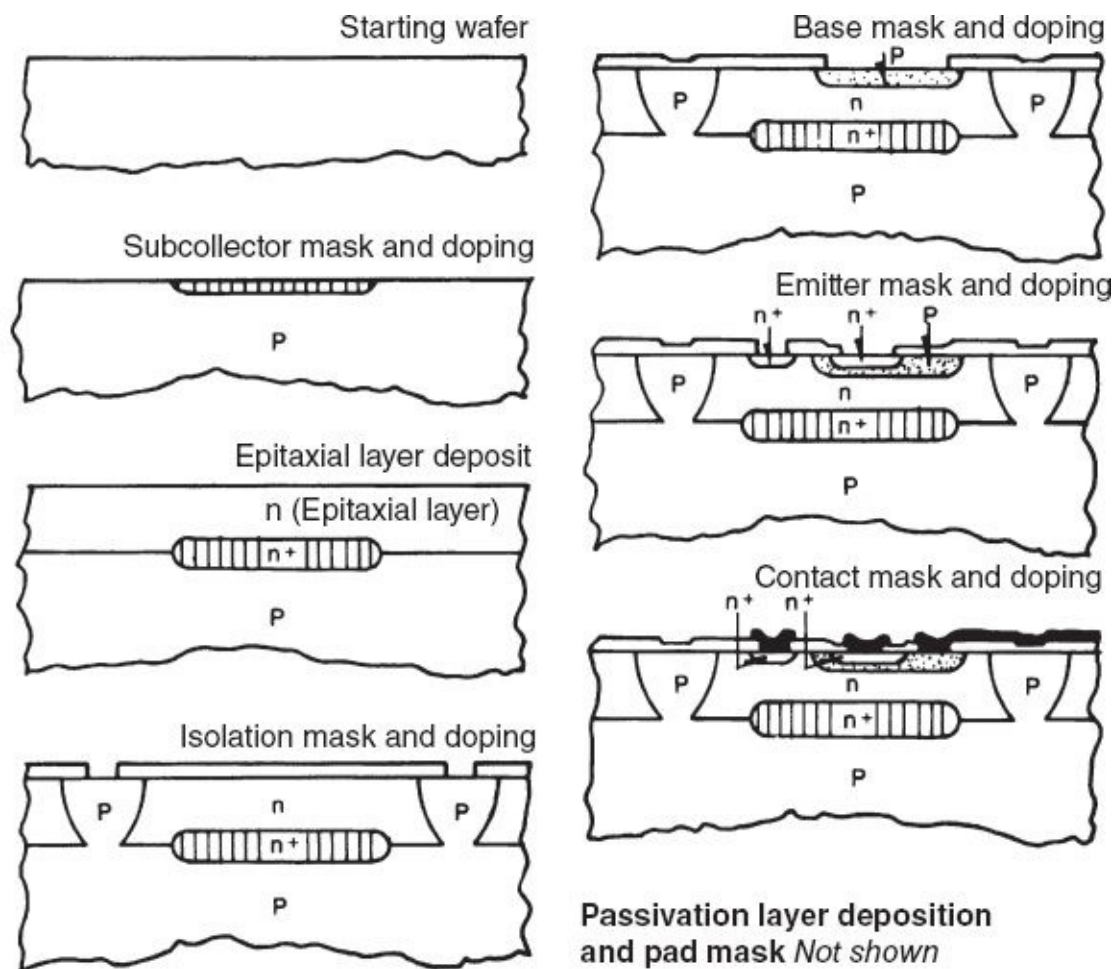


FIGURE 16.33 NPN bipolar process.

The process starts with a P-type wafer into which an N-type diffusion is made (the flow diagram does not show the oxidation and masking steps required to create the diffused layer). After the diffusion step, an N-type epitaxial layer is deposited, leaving the N-type diffused region “buried” under the epitaxial layer. The N-type region is known as a *buried layer* or as the subcollector of the transistor. Its function is to provide a lower-resistance path for the collector current as it flows out of the base region on its way to the surface collector contact.

After the deposition of the epitaxial layer, it is oxidized, and a hole is opened up on each side of the buried layer. A P-type doping step is performed deep enough to reach the P-type wafer surface. The doping step divides (*isolates*) the epitaxial layer into N-type islands, each surrounded on the sides (the doped regions) and the bottom (the P-type wafer) by P-type doped regions. Components formed on the surface of each of the islands are electrically isolated from each other (Fig. 16.33). The electrical isolation occurs because the N-P junction is “wired” into the circuit to function in the reverse-bias mode; that is, no current crosses the junction. This scheme is called *junction isolation* or *doped junction isolation*.

Note that, in the bipolar cross-section (Fig. 16.34), there is a doped region under

the transistor collector contact. This doped region is put into the surface along with the emitter N-type doping. The emitter is usually designated N^+ to indicate that it is highly doped. The N^+ region under the collector contact is present to create a lower resistance between aluminum metallization and the silicon of the collector.

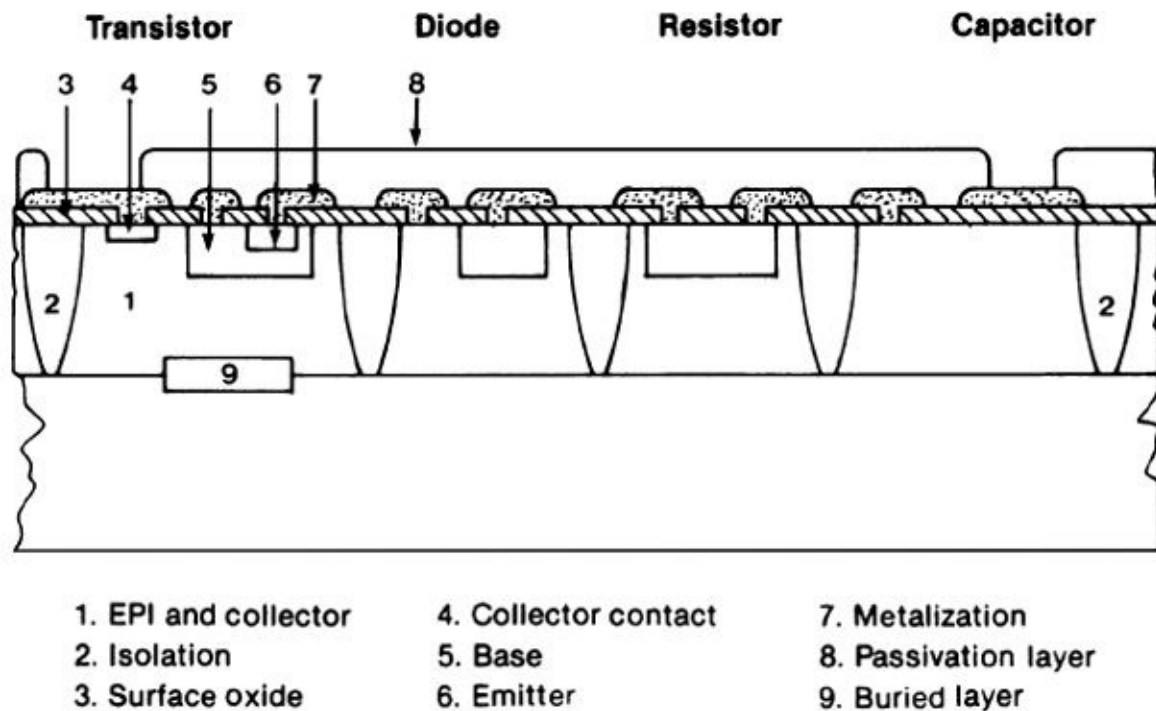


FIGURE 16.34 Bipolar structures.

Dielectric Isolation

In high-radiation environments, such as outer space or in the vicinity of atomic weapons, doped junctions produce holes and electrons that flow in the device and compromise the junction functions. Besides causing circuit component failure, the radiation swamps out the isolation protection of the doped regions. Dielectric isolation schemes provide the necessary electrical isolation and radiation protection.

The process ([Fig. 16.35](#)) starts with the etching of pockets or trenches in a wafer surface. The etching can be either an isotropic wet etch or an anisotropic dry etch. Isotropic etch profiles follow the orientation structure of the wafers. Dry-etch processes allow the shaping of the trenches. One goal of the etch step is to minimize the area of the pocket on the wafer surface. Wide pockets limit the packing density of the circuit.

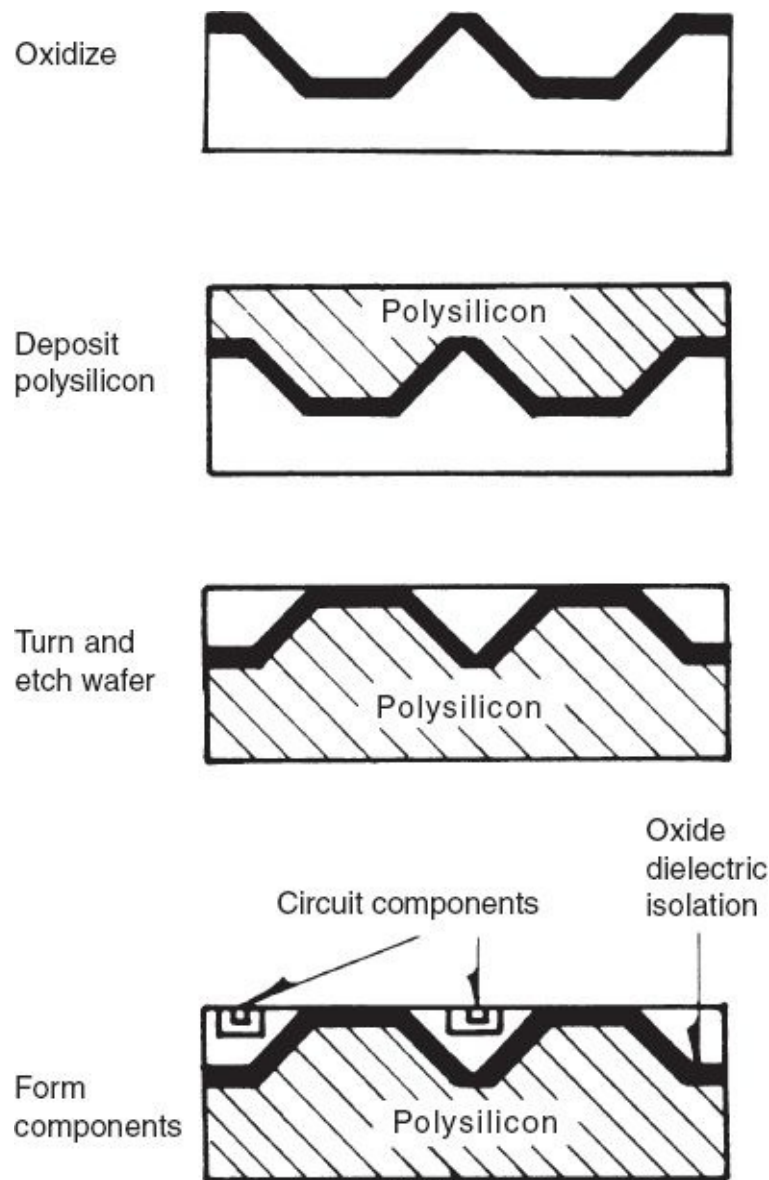


FIGURE 16.35 Dielectric isolation.

After etch, the pocket sides are oxidized and backfilled with deposited polysilicon. Next, the wafer is turned over, and the silicon of the wafer is lapped until the oxide layer is reached. These steps leave a wafer surface containing oxide-isolated pockets of the original single silicon material. The circuit components are fabricated in the silicon pockets, with each pocket being isolated on three sides by the layer of silicon dioxide. The dielectric property of the silicon dioxide prevents leakage currents in both normal and radiation environments.

Localized Oxidation of Silicon

Junction isolation takes up valuable surface real estate, and dielectric isolation is area consuming and requires extra processing steps. A popular alternative is LOCOS (Fig. 16.36). The process starts with a layer of silicon nitride deposited and etched on the wafer surface. Active devices will be formed in the area defined by the silicon nitride layer. In the partially recessed version, an oxidation follows. Oxygen will not penetrate the silicon nitride to cause the oxide to grow on the exposed silicon surface. The silicon for the silicon dioxide comes from the wafer surface and, because silicon dioxide is less dense than silicon, the oxide layer forms slightly above the original silicon surface. It is *partially recessed* relative to the wafer surface. After the oxidation, the silicon nitride is removed, leaving the area free for the formation of circuit components. A variation of the LOCOS isolation process is illustrated in Fig. 16.37b. In this process, the silicon surface is etched *before* the oxidation. By calculating the proper removal amount, the subsequent oxidized layer is *fully recessed* to the original surface. Bipolar schemes using LOCOS isolation are shown in Fig. 16.37.

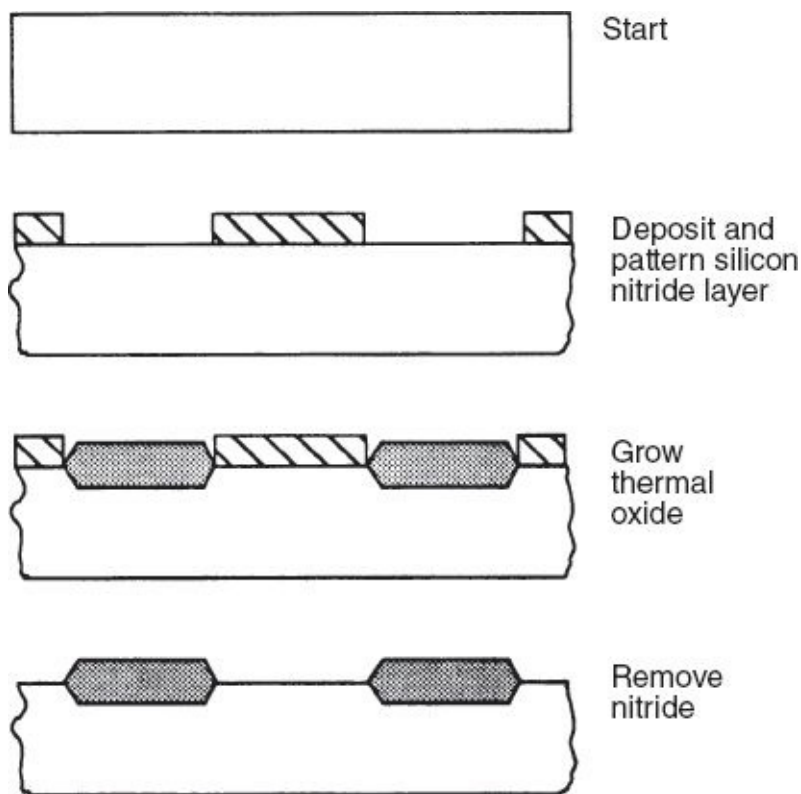


FIGURE 16.36 LOCOS process.

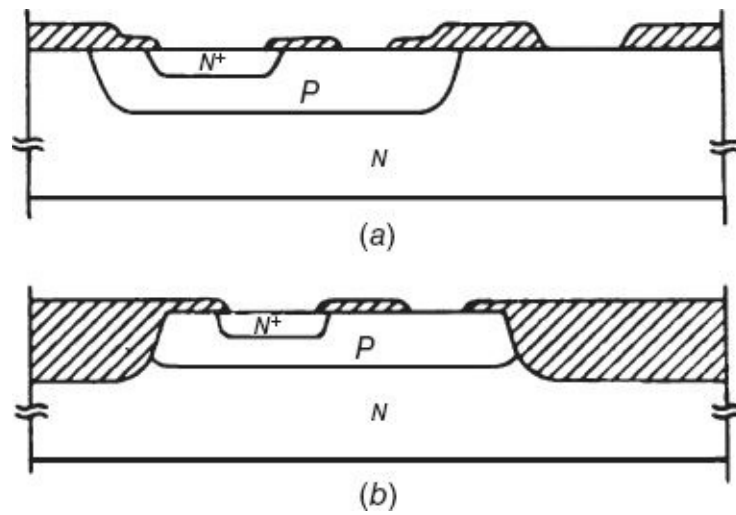


FIGURE 16.37 (a) Conventional bipolar transistor and (b) LOCOS isolated bipolar transistor. (Source: VLSI Fabrication Principles, *Ghandhi*.)

MOS Integrated Circuit Formation

MOS LOCOS Isolation

While MOS transistors are somewhat self-isolated (because they do not share common electrical parts), there is some electrical leakage between devices, especially as the spacing becomes very close. Isolation is necessary to block leakage currents. The structures are generically called *channel stops*.

LOCOS is the preferred isolation technique. However, there are several problems to overcome to make a LOCOS structure effective in advanced circuits.⁹ One problem is the *bird's beak* spur that grows under the edge of the blocking silicon nitride layer (Fig. 16.38). The beak takes up real estate, effectively enlarging the circuit. At a performance level, it induces stress damage in the silicon during the oxidation step. The stress comes from the mismatch in thermal expansion properties between silicon nitride and silicon. A solution for the stress problem is to grow a thin oxide layer under the silicon nitride film. This is called a *pad oxide*.

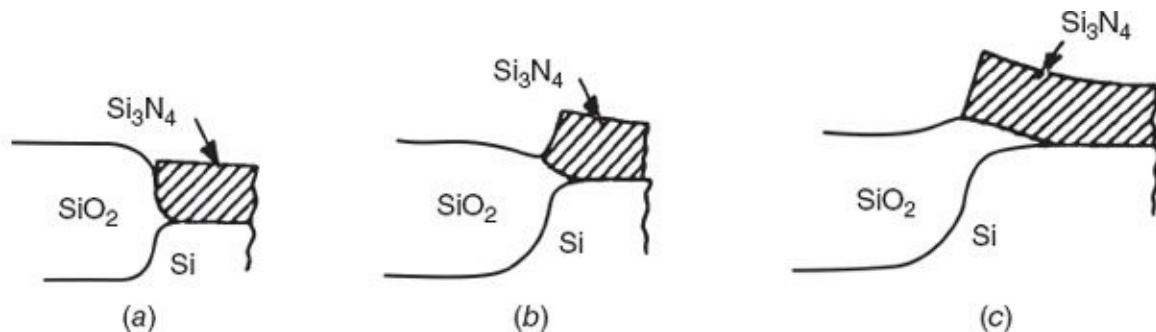


FIGURE 16.38 Bird's beak growth. (a) No pre-etch, (b) 1000-Å pre-etch, and (c) 2000-Å pre-etch. (Source: VLSI Fabrication Principles, Ghandhi.)

Minimizing the bird's beak and reducing the stress in the active device region has spurred a number of variations on the LOCOS process. One, called SWAMI, was developed by Hewlett Packard (Fig. 16.39).¹⁰ The process starts as a standard LOCOS structure. After the nitride and pad oxide, grooves (or trenches) are etched with an orientation-sensitive etchant. On $\langle 100 \rangle$ -oriented material, the groove sidewalls are at a 60° angle, which reduces stress in the silicon. Next, another stress-relieve oxide (SRO) layer is grown and covered by a layer of silicon nitride, which provides conformal coverage. An LPCVD silicon dioxide completes the sandwich before etching (Fig. 16.39c). This oxide protects the silicon nitride layer from removal. Finally, the field oxide (FOX) is grown. The length of the nitride layer governs the size of the bird's beak encroachment. Removal of the original nitride or SRO and the second nitride leaves a somewhat planarized surface for device formation. LOCOS isolation schemes usually include an ion-implanted layer

between active regions to provide further channel stop capabilities.

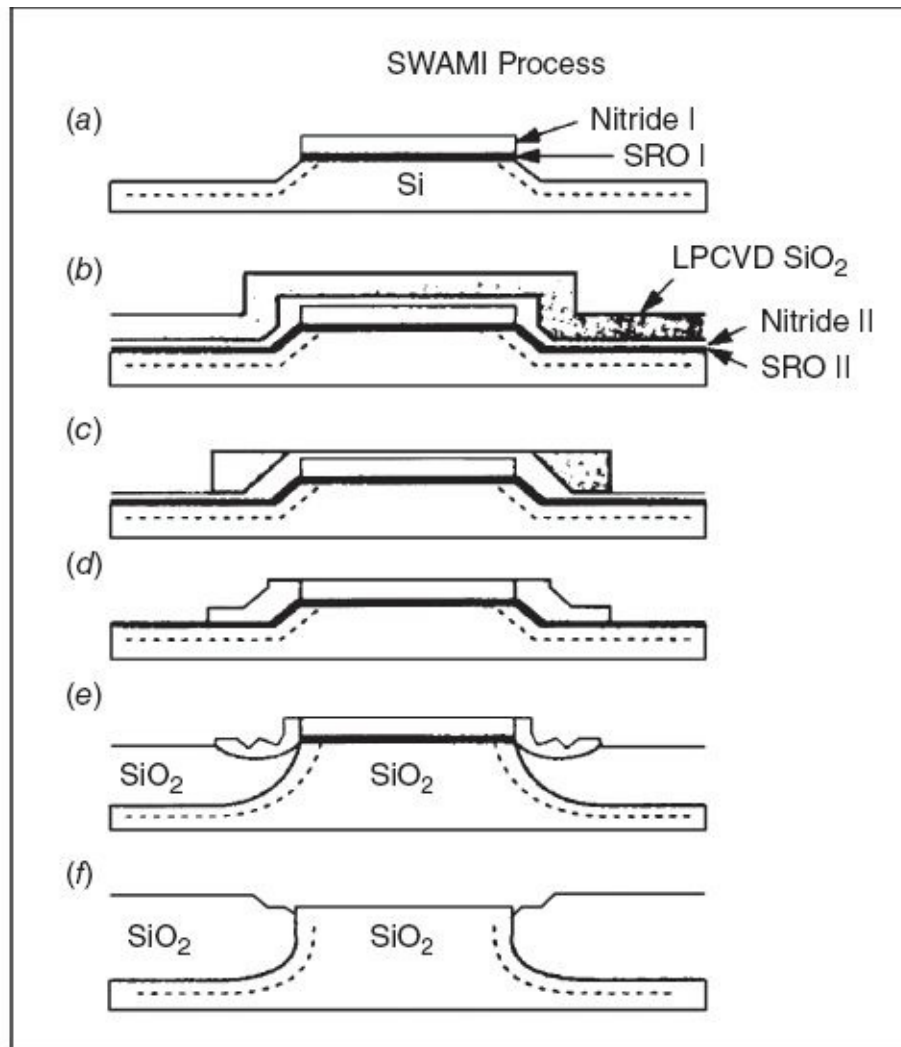


FIGURE 16.39 SWAMI process. (Courtesy of Solid State Technology.)

Trench Isolation

Trench isolation is also used for MOS circuits ([Fig. 16.40](#)). The procedure is the same as forming trench capacitors. One version, called *shallow trench* isolation, addresses the bird's beak problem that comes with standard LOCOS isolation. In this structure, a shallow trench is etched between the components and filled with a deposited dielectric ([Fig. 16.41](#)). After sidewall oxidation and dielectric fill of oxide, a CMP step replanarizes the surface.

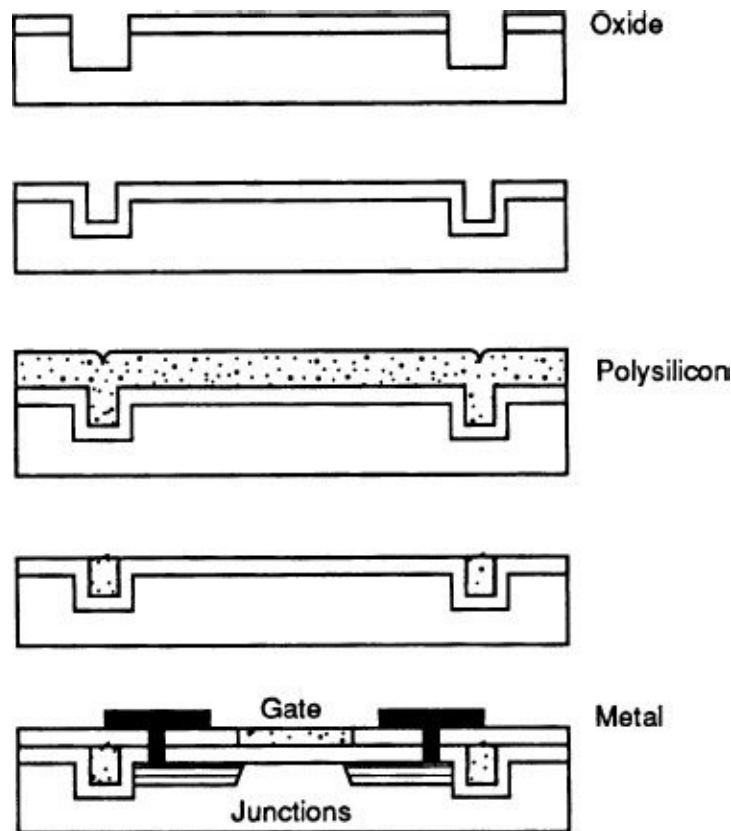


FIGURE 16.40 MOS trench isolation.

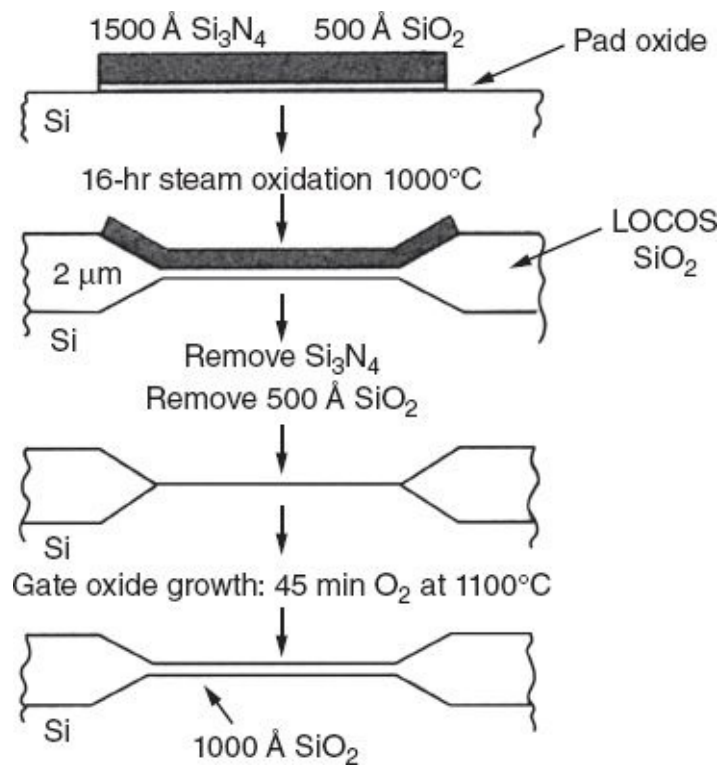


FIGURE 16.41 Shallow trench process. (Courtesy of Solid State Technology.)

CMOS

Complementary MOS (CMOS) is an MOS circuit formed with both N-channel and P-channel transistors. CMOS has become the standard circuit for many applications. It is the CMOS circuit that has made possible the digital revolution from digital watches to the development of the wireless devices that have revolutionized communications. It allows circuits on one chip that would require several chips using N-channel and P-channel only circuits. CMOS circuits also use lower amounts of power than comparable circuits.

CMOS structures ([Fig. 16.42](#)) are formed by first fabricating an N-channel MOS transistor in a deep P-type well formed in the wafer surface. After N-channel transistor formation, a P-channel transistor is fabricated. The transistor structures are a silicon gate or other advanced structures. CMOS processing uses the most advanced techniques, because smaller, more densely packed, and higher-quality components all increase the advantages inherent in the CMOS design.

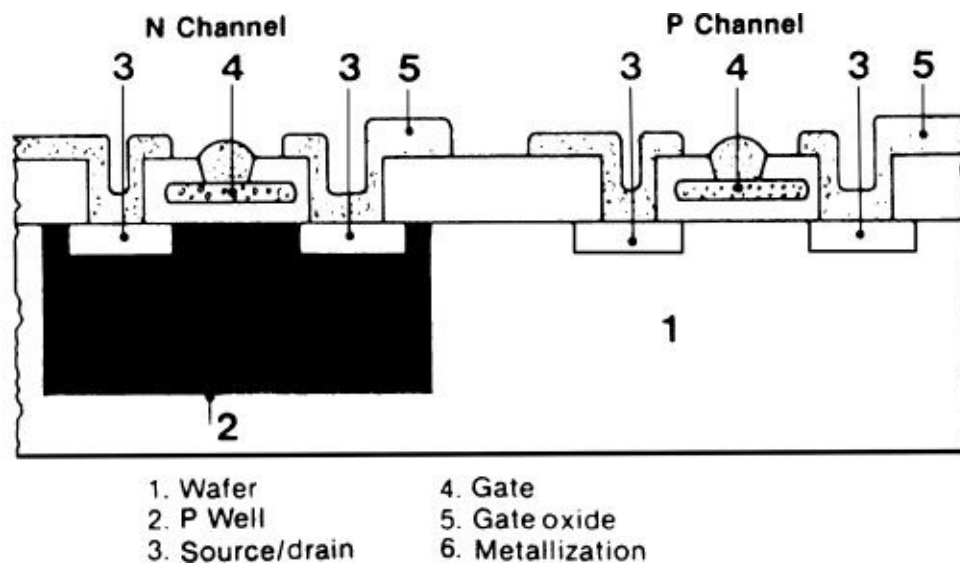


FIGURE 16.42 CMOS structure.

An isolation problem particular to CMOS structures is latch-up. [Figure 16.43](#) shows a cross-section of a portion of a CMOS chip. The side-by-side MOS transistors form a lateral bipolar (NPN) transistor. During circuit operation, the lateral bipolar transistor can function as an unwanted amplifier, increasing its output to a point where it causes the memory cell to go into a state in which it cannot be switched. This is *latch-up*. In this state, the cell cannot be addressed for its information. One solution to preventing latch-up is a low-resistivity epitaxial layer that shunts (shorts) the emitter of the bipolar transistor so that it does not “turn on,”

preventing the latch-up.

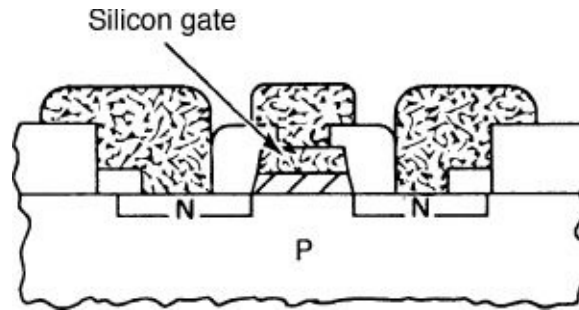


FIGURE 16.43 Cross-section of silicon gate MOS transistor.

The LOCOS process addresses latch up. Another approach is a hardened-well design using a buried layer that effectively breaks up the lateral bipolar transistor ([Fig. 16.44](#)).

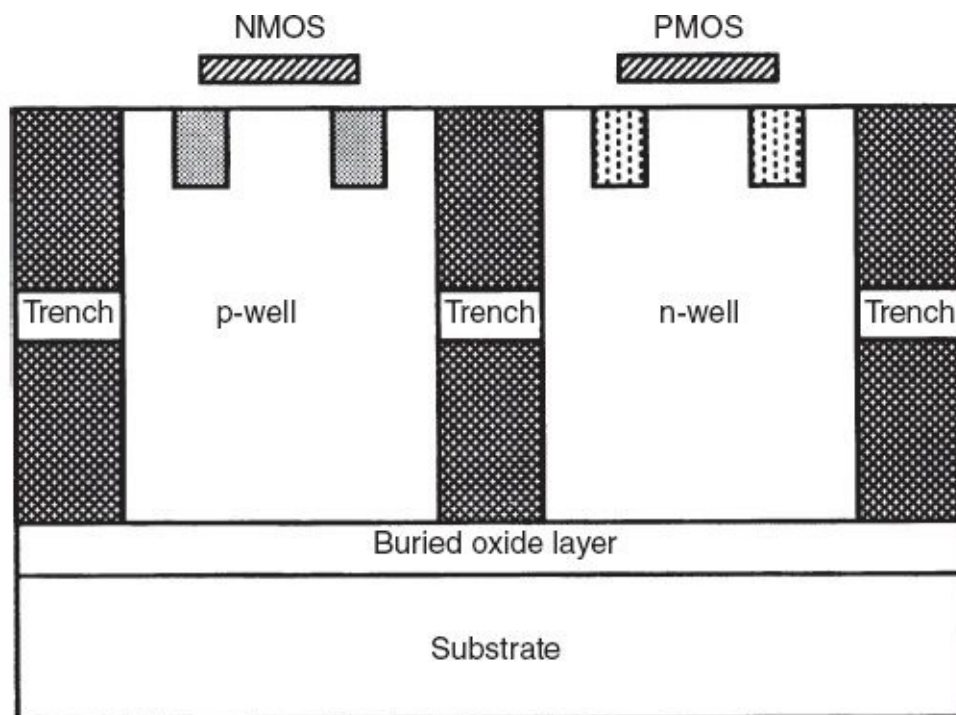


FIGURE 16.44 CMOS well structure silicon on insulator (SOI) with buried layer and trench isolation. (Source: Semiconductor International, July 1993.)

Formation of the deep P-well by diffusion techniques requires a large thermal budget that contributes to lateral diffusion problems and stress growth. Ion-implanted wells have become common. Additionally, with ion implantation, the wells can be dopant graded in the vertical direction to maximize performance. Called *retrograde wells*, they are formed by high-energy implants (MeV) such that

the bottom of the well has a higher dopant concentration. Implanted wells offer the advantage of fewer steps compared to conventional well formation^u ([Fig. 16.45](#)). Also, because the N and P wells are formed in the same surface, planarity is preserved. In conventional well processing, the two well surfaces are offset and can cause depth-of-focus problems.

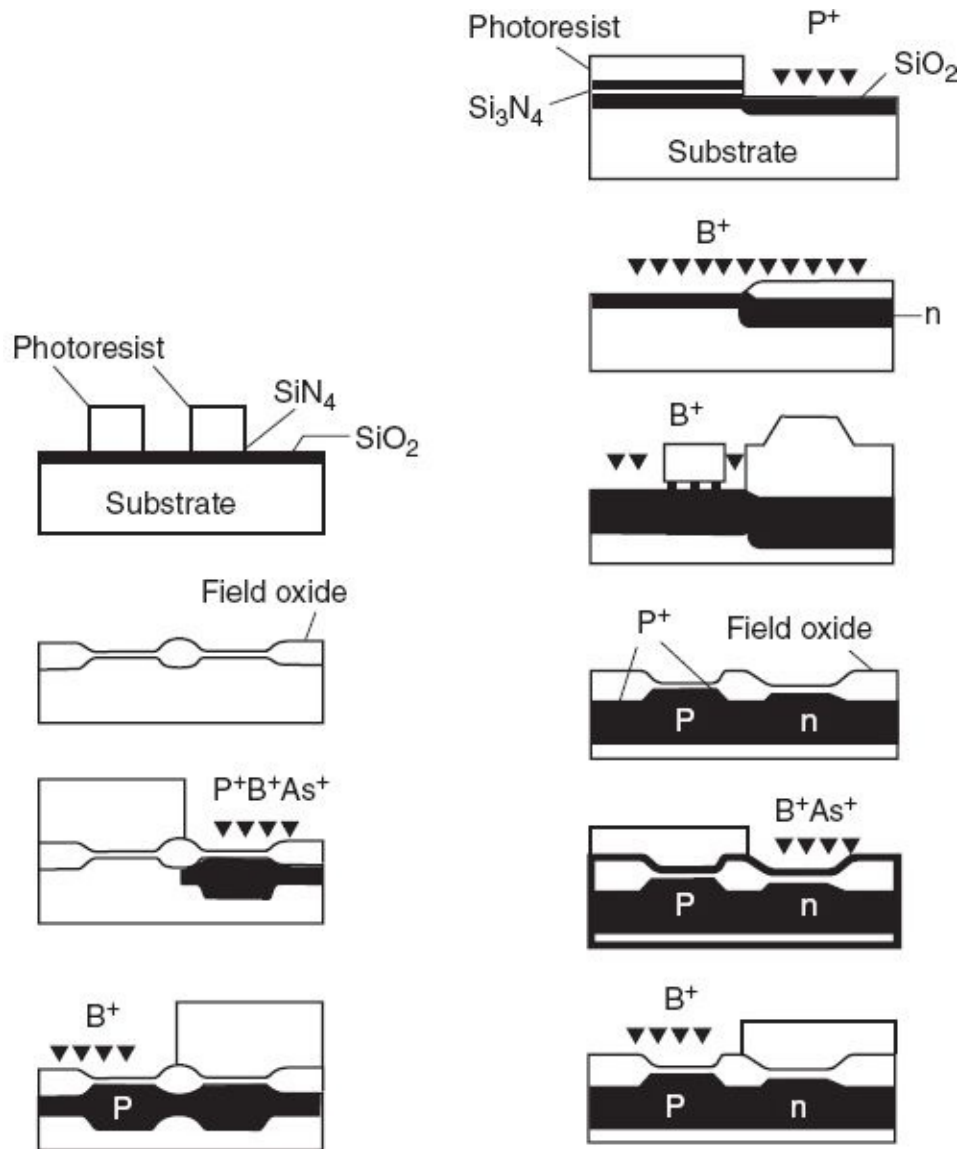


FIGURE 16.45 Implanted and conventional wells. (Source: Semiconductor International, June 1993, p. 84.)

Bi-MOS

The unique advantages of bipolar and CMOS transistors and their respective circuits come together in bi-MOS (or bi-CMOS) circuits. The circuits ([Fig. 16.46](#)) have in them bipolar, P-channel, and N-channel transistors along with memory cells (see [Chap. 17](#)). The low-power advantage of CMOS is used in logic and memory

sections of the circuit, while the high-speed performance of bipolar circuitry is used for signal processing.¹² These circuits represent great challenges to the processing area with their large die size, small component size, and large number of processing steps.

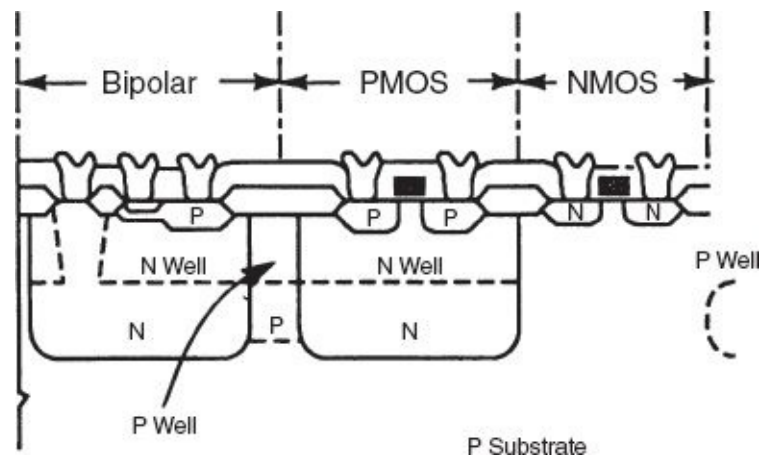


FIGURE 16.46 Bi-MOS structure.

Silicon on Insulator Isolation

Silicon on insulator (SOI) technology is the forming of CMOS and bi-CMOS circuits in thin epitaxial silicon films deposited over an insulating substrate (see [Chap. 12](#)). Elimination of a conductive substrate minimizes leakage and latch-up problems.¹³

System on (a) Chip

Innovation is often driven by new needs. One current semiconductor circuit evolution owes its development to the explosion of hand-held wireless devices, for example, cell phones, smart phones, and tablets. They are called a *system on a chip* (SOC). These devices are really computers and as such require a great many functions but with less room to house them. These include processors, graphic interfaces, data management, wireless connectivity, memory, and so forth. Fortunately, advances in microchip-fabrication processes allow these “mega” chips to be fabricated.¹⁴

The reduced space requirement is also being addressed at the packaging level. Single chips can be “packaged” with little volume beyond the chip itself. And new three-dimensional packaging schemes address the space problem with a multichip in a single package ([Chap. 18](#)).

Superconductors

Much interest has been generated by developments in superconductivity materials. Superconductivity is a phenomenon that occurs in certain materials when they are cooled close to absolute zero (-273°C). In ordinary metals, current is passed along by the flow of electrons. In an at-rest (nonconducting) state, the electrons exist in orbits around the nuclei of the atoms. To become conducting, they must gain energy to overcome the internal resistance of the particular material. The energy must be continuously supplied to maintain the current flow.

In a superconducting material, electrons exist in a “conducting state” and can support an electrical current with little or no additional energy input. The prospect of a resistance-less material has the potential of revolutionizing electronic devices.

Semiconductor researchers have investigated the superconducting effect for years. In 1962, B. D. Josephson described an effect that is now named after him (Josephson effect). When a thin oxide separates two superconducting materials, electrons will pass through the oxide with zero resistance. The structure is called a *Josephson junction*. This effect (called *tunneling*) is very complicated and requires quantum physics concepts to understand (well beyond the scope of this text). The result is that the oxide has the functional aspects of a gate, and Josephson junctions can perform basic switching, logic, and memory functions.^{[15](#)}

Microelectromechanical Systems

Semiconductor microchip-fabrication processing has given rise to a new line of products. These are microelectromechanical systems (MEMS). Generally speaking they are mechanical devices manufactured with semiconductor processes.¹⁶ Current products include microsensors and microactuators. These are basically transducers that convert physical inputs to electronic signals. Temperature, pressure, inertial forces, chemical species, magnetic fields, and radiation are examples. MEMS devices made with wafer-fabrication techniques include motors, motion sensors (as in air bag sensors), spectrometers, and miniature scanning tunneling microscopes. The future of MEMS technology is in the nano range. This would bring the devices into the atomic or molecule regime. Coupled with microelectronics, MEMS and/or integrating MEMS devices in and on integrated circuits is creating a whole new industry.

Strain Gauges

A device that takes advantage of a junction reaction to physical pressure is the strain gauge. The gauge is made by back-etching through the wafer until only a very thin membrane is left. A junction and supporting circuitry are formed in the membrane. When the membrane is deflected by some force, as in a weight scale or pressurized gas line, the semiconductor circuit produces an output proportional to the deflection. The output of the circuit is correlated with the amount of pressure on the membrane and displayed on the appropriate meter.

Batteries

Thin-film rechargeable batteries have been made from lithium and vanadium pentoxide layers. They could be formed on a package and provide backup for CMOS devices in the event of power failures.¹⁷

Light-Emitting Diodes

One effect when certain compound junctions are reversed-biased is the production of photons. Photons are a form of radiation that humans see as light. The devices are called *light-emitting diodes* (LEDs). These are the displays that are used in consumer electronics equipment and automobile displays.

The devices ([Fig. 16.47](#)) are made on gallium-arsenic-phosphide (GaAsP) wafers that are covered with thousands of diodes, wired so that they can be turned on or off individually. Groups of diodes are turned on in groups to form letters and numbers. GaAsP material produces the familiar red displays. A host of other colors are produced from different III–V and II–VI semiconductor materials ([Chap. 2](#)).

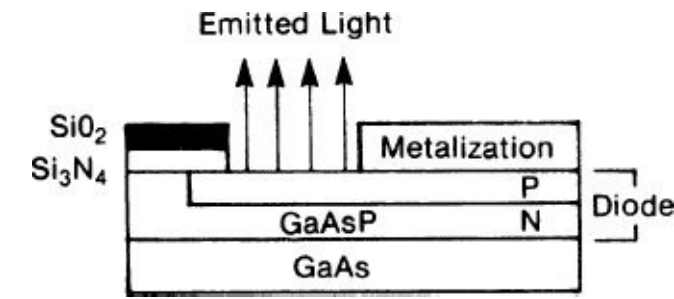


FIGURE 16.47 LED structure.

Optoelectronics

On the other side of LEDs are chips that react to light. One use sought for years is integrated opto/electronics connection in local area networks (LANs). In this use, the opto devices (sensitive to laser and other lights) are connected to a fiber-optic network or waveguide. Like the LEDs, the receivers are mostly III–V based ICs. An onboard conventional IC processes the data and sends it on through an output LED or laser.^{[18](#)}

Solar Cells

Not only can semiconductor junctions emit light, they respond to it. This property is taken advantage of in the solar cell ([Fig. 16.48](#)). The cell is composed of diodes formed in a thin layer of semiconducting material such as amorphous silicon. When sunlight strikes the junction region, a current passes across it. The current is captured by onboard circuitry.

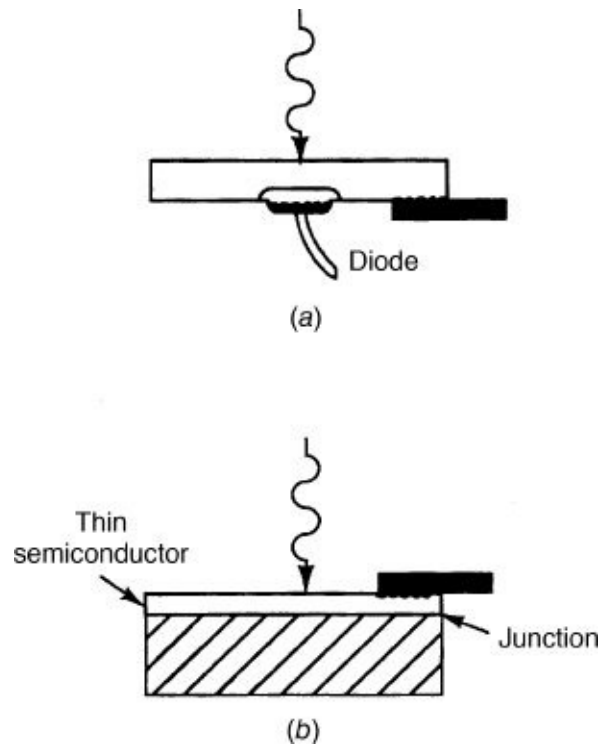


FIGURE 16.48 Light-sensitive semiconductor structures: (a) photodiode and (b) solar cell.

Temperature Sensing

Semiconductor junctions are temperature-sensitive. Heating a device will produce more current across the junction. This effect is taken advantage of in a variety of devices, such as solid-state medical thermometers and industrial control units.

Acoustic Wave Devices

Acoustic wave devices ([Fig. 16.49](#)) are nonsilicon solid-state components used in microwave communications systems. They serve the function of converting microwaves into electrical impulses. A compound material such as $\text{Be}_{12}\text{GeO}_{20}$ has the property of reacting to the wave and setting up an electrical response in a solid-state circuit formed on the chip by normal semiconductor processes.

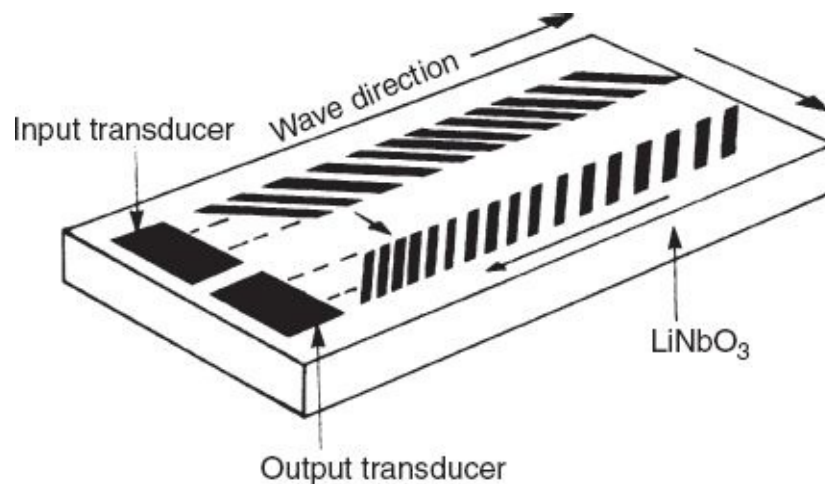


FIGURE 16.49 Acoustic wave device.

Keeping track and making sense of this structure and process complexity requires a good grounding in the basic electrical operations of the individual components and basic processes.

Review Topics

Upon completion of this chapter, you should be able to:

1. Sketch and identify the structural parts of the individual components of an integrated circuit.
 2. Explain the role and different isolation structures used for integrated circuits.
 3. Sketch and identify the operation of a bipolar and MOS transistor.
 4. List the types and advantages of the different MOS gate structures.
 5. Sketch and identify the parts of a bi-MOS circuit.
-

References

1. Camenzind, H. R., *Electronic Integrated Systems Design*, 1972, Van Nostrand Reinhold, Princeton, NJ:85.
2. Singer, P., “Gearing Up for Gigabits,” *Semiconductor International*, Nov. 1994:34.
3. De Ornellas, S., “Plasma Etch of Ferroelectric Capacitors in FeRAMs and DRAMs,” *Semiconductor International*, Sep. 1997:103.
4. Camenzind, H. R., *Electronic Integrated Systems Design*, 1972, Van Nostrand Reinhold, Princeton, NJ:141.
5. Cleavelin, C., Columbo, L., Nimi, H., et al., *Oxidation and Gate Dielectrics, Handbook of Semiconductor Manufacturing Technology*, 2008, CRC Press, New York, NY, 9–29.
6. Pauleau, Y., “Interconnect Materials for VLSI Circuits,” *Solid State Technology*, Apr. 1987:157.
7. “Industry News,” *Semiconductor International*, Cahners Publishing, Apr. 1994:16.
8. Frank, D., Hoffman, T., Nguyen, B. Y., et al., *Comparison Study of FinFETs: SOI vs. Bulk*, SOI Industry Consortium, www.soiconsortium.org.
9. Ghandhi, S. K., *VLSI Fabrication Principles*, 1994, John Wiley & Sons, Inc., New York, NY:717.
10. Wolf, S., “A Review of IC Isolation Technologies—Part 8,” *Solid State Technology*, PennWell Publishing, Jun. 1993:97.
11. Peters, L., “High Hopes for High Energy Ion Implantation,” *Semiconductor International*, Cahners Publishing, Jun. 1993:84.
12. Yarling, C. B., “M. I. Current, Ion Implantation for the Challenges of ULSI and 200-mm Wafer Production,” *Microelectronic Manufacturing and Testing*, Mar. 1988:15.
13. Yallup, K., “SOI Provide Total Dielectric Isolation,” *Semiconductor International*, Cahners Publishing, Jul. 1993:134.
14. Cunningham, A., “The PC inside Your Phone: A Guide to the System-on-Chip,” *Arstechnica*, Apr. 10, 2013.
15. Anderson, P. W., “Electronic and Superconductors,” in E. Ante’bi (ed.), *The Electronic Epoch*, Van Nostrand Reinhold, Princeton, NJ:66.
16. Gabriel, K., “Engineering Microscopic Machines,” *Scientific American*, Sep. 1995:150.
17. Bates, J., “Rechargeable Thin-Film Lithium Microbatteries,” *Solid State*

Technology, PennWell Publishing, Jul. 1993:59.

[18.](#) Singer, P., “The Optoelectronics Industry: Has It Seen the Light?”
Semiconductor International, Cahners Publishing, Jul. 1993:70.

Introduction to Integrated Circuits

Introduction

In this chapter, the general integrated circuit families and their different functions are explained. The primary products of the semiconductor industry are integrated circuits. Countless numbers and types of circuits can be created using the processes described in this text. A circuit catalog from a major integrated circuit (IC) producer such as National Semiconductor or Motorola is as large as a New York City phone book. IBM estimated that their internal circuit catalog lists over 50,000 separate circuits.

Becoming familiar with integrated circuits is not as awesome a task as these high numbers might imply. The burden is eased by the fact that most circuits fall into three primary families: logic, memory, and microprocessors that contain both logic and memory ([Fig. 17.1](#)). Within each circuit family, there are a few principal designs and functions. The multiplicity of circuits comes from the many parameter variations required for specific uses.

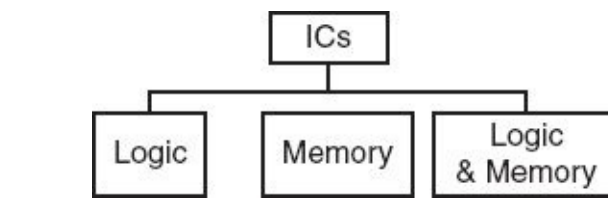


FIGURE 17.1 IC functions.

The major functional circuit categories and their circuit designs are explained in this chapter. It also looks at the future of IC circuitry from the perspective of the industry today. What the circuits will actually be like in the future can only be imagined, just as, in the 1950s, no one predicted the megabit RAM or the microprocessor.

Circuit Basics

The question of how an integrated circuit actually works is the subject of other texts. But there are some basics. All circuits are based on the processing of data in binary notation. The binary system is a way of representing any number with just two digits—for example, a zero and a one. It is actually an accounting system that keeps track of the place and value of the components of a number. Numbers can be expressed as

$$1 = 1 + 1$$

$$3 = 2 + 1$$

$$7 = 4 + 2 + 1$$

the sums of numbers. For example, $10 = 8 + 2$

Another way to express numbers is as powers of their factors. Yet another way is to express numbers as powers of their roots.

The basis of binary notation is the powers of the number 2. [Figure 17.2](#) shows the numbers 1, 2, 4, 8 expressed as powers of 2. We could also represent those numbers just as the power of 2. Thus, 1 becomes 0, 2 becomes 1, 4 becomes 2, and 8 becomes 3. Now, for the clever part, any number can be expressed as the sum of some numbers that are powers of 2. The number 25 is the sum of $16(2^4) + 8(2^3) + 1(2^0)$. In the number 25, there is one 2^4 , one 2^3 , zero 2^2 , and one 2^0 . Or the number 25 can be represented by the code 1101, each of the digits representing the presence or absence of a particular power of 2. The chart in [Fig. 17.3](#) lists some numbers in binary notation.

$1 = 2^0$
$2 = 2^1$
$4 = 2^2$
$8 = 2^3$

FIGURE 17.2 Powers of 2.

Standard Power of 2	32 2^5	16 2^4	8 2^3	4 2^2	2 2^1	1 2^0
Number						
Standard = Binary number						
1	0	0	0	0	0	1
7	0	0	0	1	1	1
18	0	1	0	0	1	0
33	1	0	0	0	0	1

FIGURE 17.3 Binary notation.

Translating numbers into binary notation is easily accomplished by establishing a grid with each column representing a power of 2. The actual number is represented by a string of zeros or ones that indicate the presence or absence of the particular powers of 2 that make up the number.

Binary notation has been known for centuries. Buckminster Fuller, in his book, *Synergistics*, has an amusing account of the use of binary coding by the ancient Phoenicians to keep track of cargo amounts. He claims that the Phoenician sailors

were considered stupid because they could not count in the system of the day, when actually they were accurately keeping track of large amounts of cargo with only two “numbers”; they were counting in binary notation.

With binary notation, only two numbers are necessary—a one (1) and a zero (0). In the discussion above, binary notation was represented by the numbers zero and one. In the physical world, binary numbers can be represented by any system that has two conditions. [Figure 17.4](#) shows several different ways to code the number 7. The last row could represent binary coding by the off-on states of a transistor or memory cell.

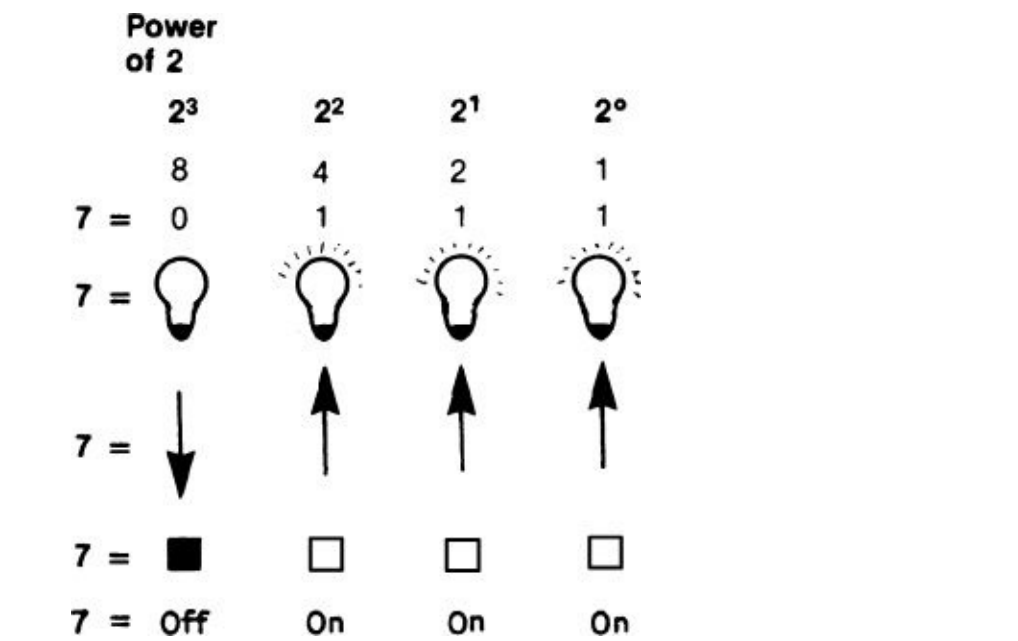


FIGURE 17.4 Binary representations of the number 7.

Inside a circuit, numbers are coded, stored, and manipulated in their binary code. The numbers can be stored and manipulated because capacitors can be charged to have a charge or not have a charge, and transistors can be either on or off. The smallest piece of information in a circuit is called a “binary digit” or “bit.” The binary coding system is simple. The problem of how the coded numbers could be added, subtracted, and multiplied was solved by George Boole, a nineteenth-century mathematician. He developed a logic system capable of handling numbers in binary notation. Until the development of computer logic, his Boolean logic (or Boolean algebra) was an academic curiosity.

Specific numbers in computer systems are processed as binary (zeros and one) and called bits or words. Chips and computers are designed to handle a bit of a specific size. An eight-bit machine manipulates numbers with eight binary bits at a time. A 64-bit machine can handle a number composed of 64 binary bits. The more bits a machine can handle at one time, the faster and more powerfully it processes

data. Every eight bits is known as a byte. Thus, an 8 MB (megabyte) storage capacity can hold 8,000,000 bits of information. And there is a code for identifying bit capacity for circuits and processing levels ([Fig. 17.5](#)).

Decimal		Metric
1000	k	<u>kilo</u>
1000 ²	M	<u>mega</u>
1000 ³	G	<u>giga</u>
1000 ⁴	T	<u>tera</u>
1000 ⁵	P	<u>peta</u>
1000 ⁶	E	<u>exa</u>
1000 ⁷	Z	<u>zetta</u>
1000 ⁸	Y	<u>yotta</u>

FIGURE 17.5 Table of “bit” values.

Integrated Circuit Types

A solid-state integrated circuit is composed of a number of separate functional areas. Each chip, regardless of the circuit function, has an input and encode section where the incoming signals are “coded” into a form that the circuit can understand. The majority of the circuit area contains the circuitry required to perform the circuit function, either memory or logic. After the data is manipulated by the circuit, it goes to a decode section where it is changed back into a form that is usable by the machine’s output mechanism. The circuit’s output section sends the data to the outside world. These are identified as the I/O (In/Out) sections of the circuit.

Although, this is an overly simplified explanation of a circuit, it illustrates the fact that the interior of a chip is composed of definite separate functional areas. In many circuits, these areas perform the same functions as the main parts of a computer.

Circuit types fall into three broad categories: logic, memory, and logic and memory (microprocessors). *Logic circuits* perform a specified logical operation on the incoming data. For example, pushing the plus (+) key on a calculator instructs the logic portion of the chip to add the numbers presented to it. An on-board automobile computer goes through a logical operation to direct the signal from a sensor indicating an open door to light up the correct warning light on the dashboard.

Memory circuits are designed to store and give back data in the same form in which it is entered. Pushing the pi (π) key on a calculator activates the memory part of the circuit where the value of π (3.14) is stored. The value 3.14 is displayed on the screen. Every time that key is activated, that value is displayed. Of course modern memory circuits store and manipulate huge amounts of data, well into the terabit

arena.

Microprocessors are a third circuit type that combines both logic, arithmetic functions, and memory. In 1972, Intel Corporation introduced the first practical microprocessor. It was the microprocessor that allowed the design of powerful personal computers, digital watches, and one-chip calculators and the transfer of so many business machines to solid-state electronics, from phone systems to vending machines. Microprocessors can be programmed to perform many different circuit functions. To accomplish this, they contain logic and memory circuitry as well as the necessary encode, decode, input, and output sections. The microprocessor has been dubbed “a computer on a chip.” While it contains all the functional areas of a computer, it is not truly a complete computer. Even simple computers require vast amounts of memory capacity that microprocessors do not have. Within many personal computers, microprocessors function as the central processing unit (CPU). Additional memory chips have to be included to make the computer of practical use. However, System on a Chip (SoC) circuits are being introduced (See section *The Next Generation*, pg 462).

Actually, every integrated circuit contains both logic capability and memory sections. For example, the logic circuitry of a calculator must have certain constants stored in a memory section to perform calculations. And memory circuits must have some logic functions to direct the flow of electrons and holes to the right parts of the circuit for storage.

Logic Circuits

Analog-Digital Logic Circuits

Logic circuits fall into two main categories: analog and digital ([Fig. 17.6](#)). Analog logic circuits were the earliest logic circuits developed. An analog circuit has an output that is proportional to the input. Digital circuits, on the other hand, feature a predetermined output in response to a variety of inputs. A wall light dimmer is an analog device. Turning the control varies the voltage to the dimmer, which in turn varies the brightness of the light. A standard on-off light switch is a digital device. Only two brightness conditions are possible: on or off. Early audio circuits were based on analog circuits but the more precise digital circuits have taken over.

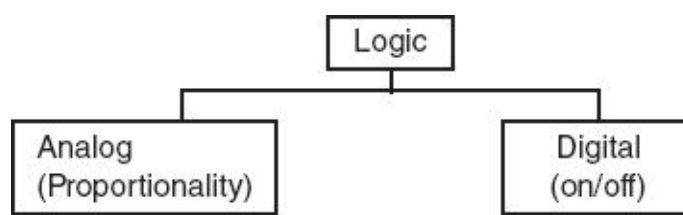


FIGURE 17.6 Logic circuit types.

Analog Logic Circuit Types

Analog circuits were the first type designed in integrated form. The home computer hobby kits of the 1950s were the analog type. These simple circuits were based on Ohm's law ($R = V/I$). The circuit contains a resistance meter and a means for generating a current and measuring a voltage. The three quantities are related by Ohm's law. Any other three variables similarly related can be represented by the resistance, voltage, and current. Varying one changes the other two. The circuit thus becomes a computer for solving any equation of the form $A = B/C$.

The accuracy of analog circuits is dependent on the precision of the relationship between the input and the output. In the simple computer illustrated, accuracy is dependent on the precision of the components in the circuit, the clarity of the meters for setting the input and reading the output, and the immunity of the circuit to outside "noise." Unless the circuit contains a section to regulate incoming voltage levels, a change in the line voltage would alter the output, and hence the accuracy.

Both simple and complex analog circuits are vulnerable to variations in the incoming signal and to internal noise. Analog circuits are also dependent on precise control of the resistor values. Unfortunately, diffused resistors cannot be fabricated with a resistance variation from design value better than 3 to 5 percent, which is unacceptable for many applications. Greater resistance precision is gained by the use of matched resistor pairs, in which the effective resistance in the circuit is the difference between two resistors. This difference can be more tightly controlled than with one resistor alone.

Ion implantation also provides the analog circuit designer with a tool for producing resistors with a higher degree of control. Many analog circuits feature thin-film resistors to achieve the required precision. The growth and popularity of digital circuits is based on their ability to produce a set output every time, regardless of noise or signal variations in the circuit. If a 5-V output is required, the digital circuit will deliver exactly that level.

Digital circuits, however, do not respond as fast as linear circuits. The term used in electronics is *real-time response*. In some applications, such as airplane controls, real-time response is mandatory. Increased digital circuit speed is facilitating the encroachment of digital circuitry into this traditional use of analog circuits. A major advantage of digital circuits over analog circuits is in generalpurpose computers. Analog circuits are more difficult to design to respond to a general range of problems. All modern generalpurpose computers are based on digital circuits.

The most popular use for analog circuits is in amplifiers. They are designed in a variety of configurations, for many different applications. All have the same basic principle—the incoming signal or pulse is amplified. Audio circuits require amplification of a weak signal from an input (radio signal, CD, etc.) so as to

produce the level required to operate a speaker.

The real-time aspect of analog circuits also makes them real-world circuits. Wherever there is a real-world measurement, such as temperature or movement, analog circuits are used. Even when the majority of the circuitry in a system is digital, analog circuits are often part of the interface with the outside world.

Most analog amplifier circuits are of the differential operational type. These circuits produce an output voltage amplified from and proportional to the difference of two input signals. Bipolar technology is favored for these circuits, because bipolar circuits are modular electric current devices and are better suited to the applications required of analog circuits.

The output signal of an analog device can have a “one-to-one” relationship with the input signal. These circuits are called *linear*. If the input is changed, the output changes in a linear manner. So many analog circuits are of linear design that the two terms are often used interchangeably. However, there are nonlinear circuits, such as those featuring a logarithmic relationship between the input and output.

Logic circuitry is built around the logic gate. A gate controls and directs passage through a barrier. The size and design of a gate influence the amount of passage allowed. For example, a room with many “in” doors and only one “out” door is an example of a “gate” configuration. Many people can enter the room, but their exit is restricted, because only one door is provided. This example gate can also be operated in reverse, allowing people to enter through only one door and leave through many.

Electronic logic gates perform similar functions with electrical signals. In a circuit, they perform the necessary logic operation by the dictates of Boolean logic. A discussion of their incorporation in logic design is beyond the scope of this text. The point for this text is that gates, both analog and digital, are formed in a logic circuit by wiring together various components.

Custom-Semicustom Logic

Using any of the logic gate approaches listed, hundreds of thousands of different logic circuits can be constructed. They vary from custom small-volume circuits to off-the-shelf standards. The bulk of logic circuits require some degree of customization. Several design and fabrication approaches are used to deliver custom and semicustom circuits to the customer at reasonable costs.

The approaches are: 1. Full custom 2. Standard circuit—custom gate pattern 3. Standard circuit—selective wiring gate 4. Programmable array logic **Full Custom**

A full custom-designed logic circuit is specified by the customer, who pays for design and mask-making fees along with the fabrication costs. This approach is expensive and lengthy and is not geared for experimenting with different circuits in the design stage of a project. Custom-designed circuits are not cost effective in

quantities of less than 100,000.

Standard Circuit—Custom Gate Pattern

This fabrication process starts with a standard logic circuit design, but only the gates required for the particular application are formed during the fabrication process. The input, output, and other circuit sections are standard for a family of circuits.

Standard Circuit—Selective Wiring Gate Arrays

This system is similar to the custom gate approach but is based on a standard circuit design for most of the fabrication process. These circuits are built with a standard number of gates. This gate section is called the *array*, and the circuit is known as a *gate array*. Working with the basic design, customers can instruct the fabrication department to wire together only the gates required to produce their custom circuit logic function.

The result is faster turnaround time than full custom processes can achieve, and moderate cost. The cost per logic function of gate arrays is higher than that of a custom circuit produced in production quantities. The larger gate section required to allow many different circuit variations results in a larger chip. This larger chip size leads to a higher manufacturing cost per chip and/or a lower yield.

The wafers receive a common process up to the contact mask. The contact mask is customized to form contacts only to the gates required for the particular circuit. After metal deposition, only the gates with contacts are wired into the circuit. A variation of this process is to open the contacts in all the gates but use a customized metal mask that wires in only the wanted gates.

Programmable Array Logic and Field Programmable Gate Arrays

Each of the three systems described above requires the chip manufacturer to do the “customizing” before shipping. This requirement can result in delivery or scheduling problems and generally forces the user to buy a minimum quantity of parts. Monolithic Memories, Inc. (MMI) addressed this problem with the introduction of their programmable array logic (PAL) line of circuits in 1978. MMI applied the programmable fuse technology used in memory products to logic circuits. The result was a field programmable (custom) logic circuit.

More recently field programmable gate arrays (FPGAs) are providing field programmable circuits. Typically the arrays are complemented with microprocessors and other circuit sections to provide complete systems in a chip such as the Xilinx Zynq™-7000 all programmable SoC. These circuits contain unassigned logic blocks that can be custom-assembled electronically to perform the desired logic functions.

Memory Circuits

Around 1960, industry forecasters began predicting that solid-state memory circuits would overtake the traditional core memory. The advantages of solid circuits were their reliability, small size, and faster speed. This prediction was made every year until the early 1970s, when solid-state memories finally did surpass core memory. The major factor that prolonged the life of core memory was lower cost.

Metal-oxide semiconductor (MOS) is the favored transistor structure for memory circuits. During the 1960s, however, the cleanliness requirement for MOS processing was not reliably available. High-yield MOS processing also requires accurate alignment and clean thin gate oxides. These processes were not fully developed in those years. The resultant low process yields kept MOS memories more expensive than core memories.

With process improvements and improved silicon gate structures, complementary metal-oxide semiconductor (CMOS) is the memory technology of choice. Some bipolar memories are favored for their fast speed and switching capabilities. While logic circuits can be (and are) made in MOS technology, most MOS circuits produced are memories, with the majority incorporated into computers. They are also used in microprocessor-based products, which require auxiliary memory chips. There are two principal types of memory circuits: volatile and nonvolatile ([Fig. 17.7](#)).

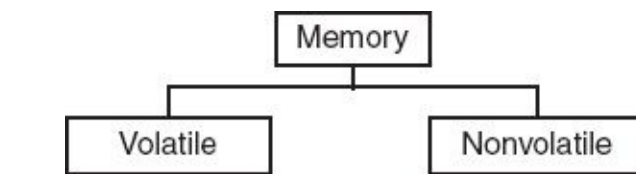


FIGURE 17.7 Memory circuit types.

Nonvolatile Memories

A nonvolatile memory device is one that retains its stored information when it loses its power. An example of this is a compact disk (CD), which is an information-storage device. If power to the record player is lost, the songs are not lost from the CD itself. A number of nonvolatile memory circuits are listed in [Fig. 17.8](#).

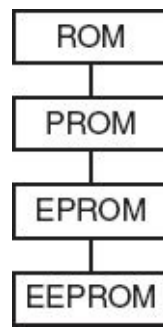


FIGURE 17.8 Nonvolatile memory.

In addition to these IC nonvolatile memory systems other examples include most types of magnetic storage devices such as hard disks, floppy disks, magnetic tape, and optical discs.

ROM

In integrated circuits, the ROM design is the principal nonvolatile circuit. ROM stands for *read-only memory*. The sole function of this type of circuit is to give back precoded information. The information required in the circuit is specifically designed into the chip memory array section during fabrication. Once the chip is made, the stored information is a permanent part of the circuit.

Other memory types have *read and write* capability. That is, they can receive and store information from an input device [keyboard, magnetic tape, compact disk (CD)]. A CD is a nonvolatile ROM device. A magnetic tape is also an example of a nonvolatile device with both read and write capabilities, because information can be erased and rerecorded.

In a calculator, the constants and the rules required for the math operations are available in the ROM section. ROM circuits, like logic circuits, number in the hundreds of thousands. Although there are many standard types, the industry also produces many custom ROM circuits. The choices offered to the user in selecting a standard or custom chip are similar to those available with logic circuits. The user can buy a standard circuit, specify a variation on a standard basic circuit, design a total custom circuit, or buy a PROM, EPROM, or EEPROM.

PROM

A programmable read-only memory (PROM) is the memory equivalent of a PAL. Each memory cell is connected into the circuit through a fuse. Users program the PROM to their own memory circuit requirements by blowing fuses at the unwanted memory cell locations. After programming, the PROM is changed to a ROM, and the information is permanently coded in the chip, where it becomes a read-only memory.

EPROM

For some applications, it is convenient to change the information stored in the ROM without having to replace the whole chip. EPROM (erasable programmable ROM) chips are designed for this use. The erasable feature is built in with the use of MMOS (memory MOS) transistors detailed in [Chap. 16](#). The transistors can be selectively charged (or programmed), and they hold the charge for a long time—which holds the information in a nonvolatile fashion. Programming is by the mechanism of hot electron injection. When reprogramming is required, the charge in the transistors can be drained off (erasing the memory) by shining ultraviolet light on the chip. Reprogramming of the chip takes place by removing it from the circuit and putting in new memory information with an external programming machine. A typical EPROM can be reprogrammed up to 10 times.

EEPROM

The next level of convenience in memory design is the ability to program and reprogram the chip while it is in a socket in the machine. This convenience is available with the EEPROM (or E² PROM), standing for *electronically erasable PROM*. Programming and erasing take place by pulses from the outside that place charges in selected memory cells or drain the charges away. Programming is by the same mechanism used for EPROMs, hot electron injection. Charge is drawn from the memory cell by a mechanism called Fowler-Nordheim tunneling. However, a larger memory cell size is required and a commensurate reduction in chip density.

Flash Memories

A flash memory is a form of EEPROM. It is a one-transistor cell design,¹ and like EPROMs can be programmed and reprogrammed while in their sockets in a system. Additionally, blocks, or the entire array, can be erased at one time increasing operational functioning.

Also flash memory is the memory part of USB storage devices based on a NAND gate configuration. This type of flash memory is best at storing and restoring large amounts of memory, thus their popularity in USB memory storage. The NOR configuration is faster at computation and finds use in computers.

Volatile Memories

Semiconductor circuit and computer design involves the constant evaluation of tradeoffs. In the case of memory, nonvolatile memory provides protection against power loss, but these memories are frequently slow and not very dense. More importantly, none of the circuits described above has a write capability, an essential feature in operating a computer. New information, such as a change in pay status, must be conveniently entered into the computer and stored temporarily while the

new check is being written. Memory must also be easily erasable so the computer can quickly process new information or accept a completely new program. Several memory circuit designs are used to produce fast and high-density memory circuits. Both are of the volatile type; that is, when power is lost to the chip, all the stored information is lost; information presented on a computer screen, and not saved, is eligible for loss if the power to the computer goes off.

RAM

One type of circuit used for high-density memory storage is random-access memory, or RAM. “Random” refers to the ability of the computer to directly retrieve any information stored in the circuit. Unlike a serial memory, the RAM design allows the chip to find the exact information asked for, wherever it is located in the computer memory. This feature allows faster information retrieval and makes the RAM the principal memory circuit in computers.

Dynamic Random-Access Memory

Dynamic random-access memories (DRAMs) come in two principal designs: dynamic and static ([Fig. 17.9](#)). A dynamic memory design called a DRAM, for dynamic RAM, is used in great quantities in computer memories. The memory cell design is based on only one transistor and a small capacitor ([Fig. 17.10](#)). The information is stored in it by a charge built up in the capacitor. Unfortunately, the charge drains away very rapidly. To combat this problem, the memory information must be re-inputted to the circuit on a constant basis. The term for this function is *refresh*. The refreshing of the circuit occurs many thousands of times per second. Dynamic RAMS are vulnerable to both power loss and interruption, and to problems with the refreshing circuit.

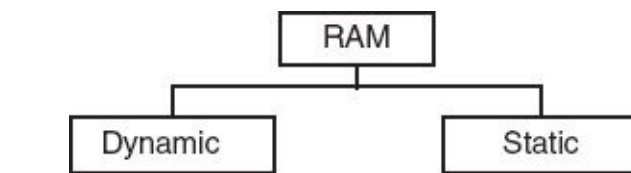


FIGURE 17.9 RAM cell designs.

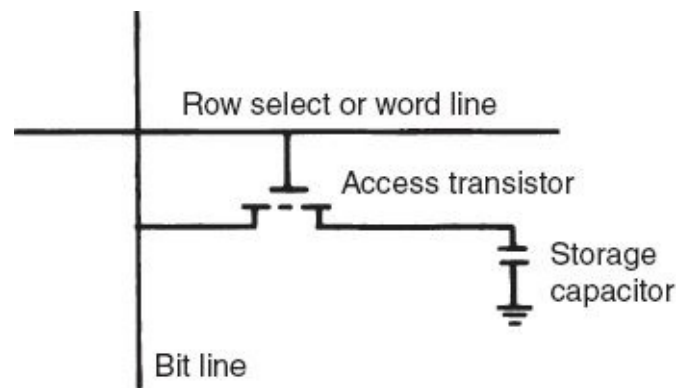


FIGURE 17.10 Dynamic RAM cell schematic.

The goal of DRAM design is a small-cell design for high-density and closely spaced components with small and thin parts for speed. These requirements have driven DRAM design and processing to the highest levels of the technology. All the advantages available by advanced, state-of-the-art equipment and processing are applied to DRAM circuits. This fact makes them the industry's leading-edge circuit.

Static Random-Access Memory

Static random-access memories (SRAMs) are based on a cell design that does not need a refresh function. Once the information is put into the chip, it will stay as long as the power remains on. This is accomplished with a cell containing several transistors and capacitors ([Fig. 17.11](#)). The information is stored as conditions of the transistors are alternately on or off. Information can be read and written with a SRAM cell much faster than with a RAM design, since transistors can be switched faster than capacitors can be charged and drained. The penalty paid for this lesser degree of volatility and speed is loss of space. The larger cell design makes static memories less dense than DRAMs.

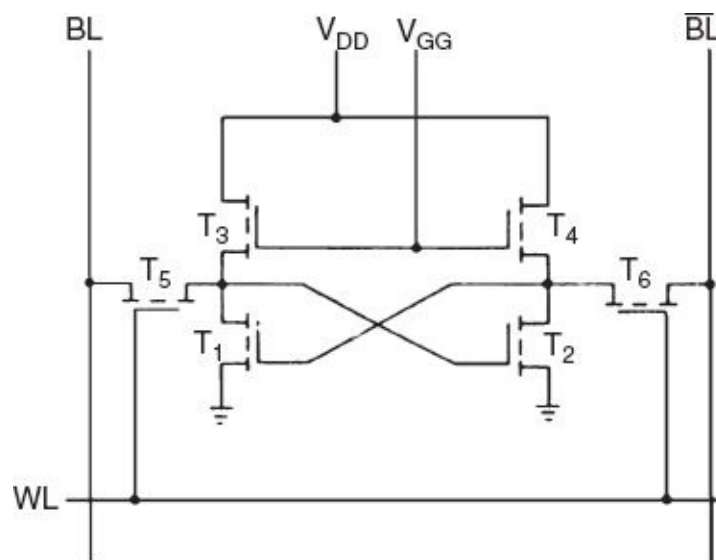


FIGURE 17.11 Static RAM cell schematic.

Memory capacity is measured by the number of bits that can be stored. A 1-k RAM has a capacity of 1024 bits of information; 1024 is a power of 2. A 64-k RAM actually has a capacity of 65,536 bits of information. RAM capacity has expanded rapidly, with larger megabit memories (64 and higher) expected to be produced in quantities with presently identified technology. Each step upward in RAM capacity places greater pressure on wafer processing and yield improvement. The nature of the semiconductor chip business is exemplified by the 64-k RAM, introduced by IBM in 1977. The chips were soon available in the merchant market, priced at over \$100 per circuit. By 1985, competition and yield improvements had lowered the price to under \$1 per circuit.

Ferroelectric Memories

Ferroelectric materials (see [Chap. 16](#)) used in capacitors create memories that operate faster than conventional SiCMOS technology capacitors.² The challenge is the integration of this nonsilicon technology into a standard silicon process.

Redundancy

Redundancy is the inclusion of extra circuit components in the design. If one or more of the components do not work, others are available that do. The tradeoff for redundancy is larger chip size. Also, extra circuitry is required within the main circuit to detect the functioning and nonfunctioning components and direct the selection of a functioning component. Although this approach to higher yield has been discussed for years, it has not yet become a mainstay of circuit design. This is due to the problem of locating the working and nonworking components and wiring the working ones into the circuit.

The Next Generation

Since the 1950s, the semiconductor industry has maintained a rather constant rate of product evolution. In general, there has been a new product generation every 3 years.³ Within each generation, the density of memory chips increased four times and logic chips two to three times. Every two generations (6 years), the feature size decreased by a factor of two, with chip area and package pin count increasing by a factor of two.

Predicting the future is always difficult, but there are identified end points for chip circuits as we now know them. Chip densities of 16 G (billion) bits for memory and 20 M (million) gates for logic are ambitious goals. The industry will have to

develop the processes and equipment for 450-mm diameter wafers. Die sizes are in the 1000-mm² range (1.2 in per side) or larger. The moves in circuits will be additional disparate IC functions combined in circuit chips. These are being driven by the explosion of handheld devices that are approaching the capabilities of laptops. The edge of this trend is the System on a Chip (SOC) evolution ([Fig. 17.12](#)).

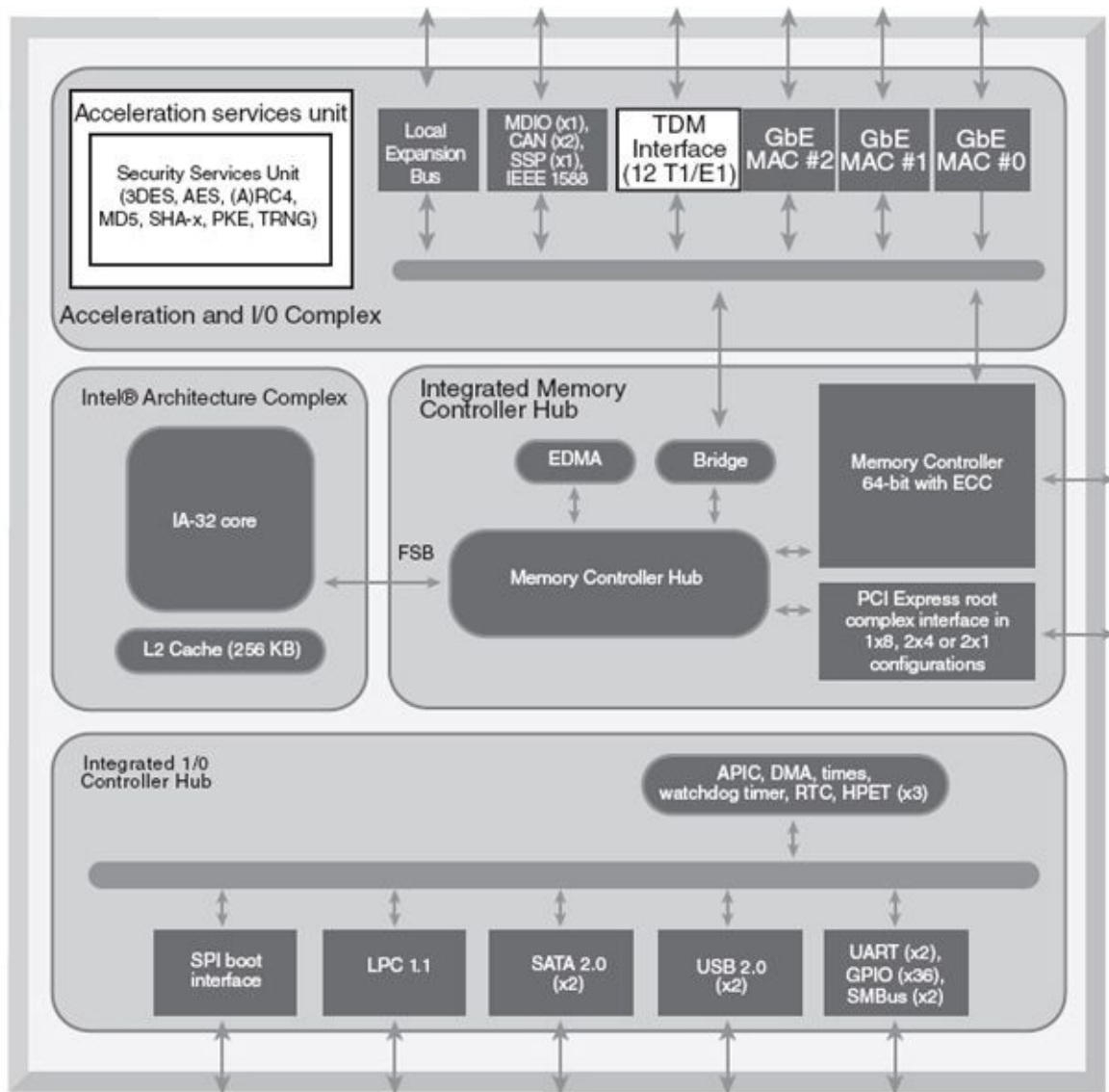


FIGURE 17.12 Intel system on a chip block (SoC) diagram. (Courtesy of Intel Corporation.)

New areas for chip development are speech recognition, expert systems, and the continuing “microchipization” of automobiles and most power-consuming products (control and conservation).

Moore’s law is approaching its limits on a flat plane, such as a wafer surface. You have heard this before and one day it will happen in the x-y plane. But Moore’s law has stayed alive in component and function density as the industry has gone up into

the z plane. This has been evident in the multimetal layer development ([Fig. 17.13](#)).

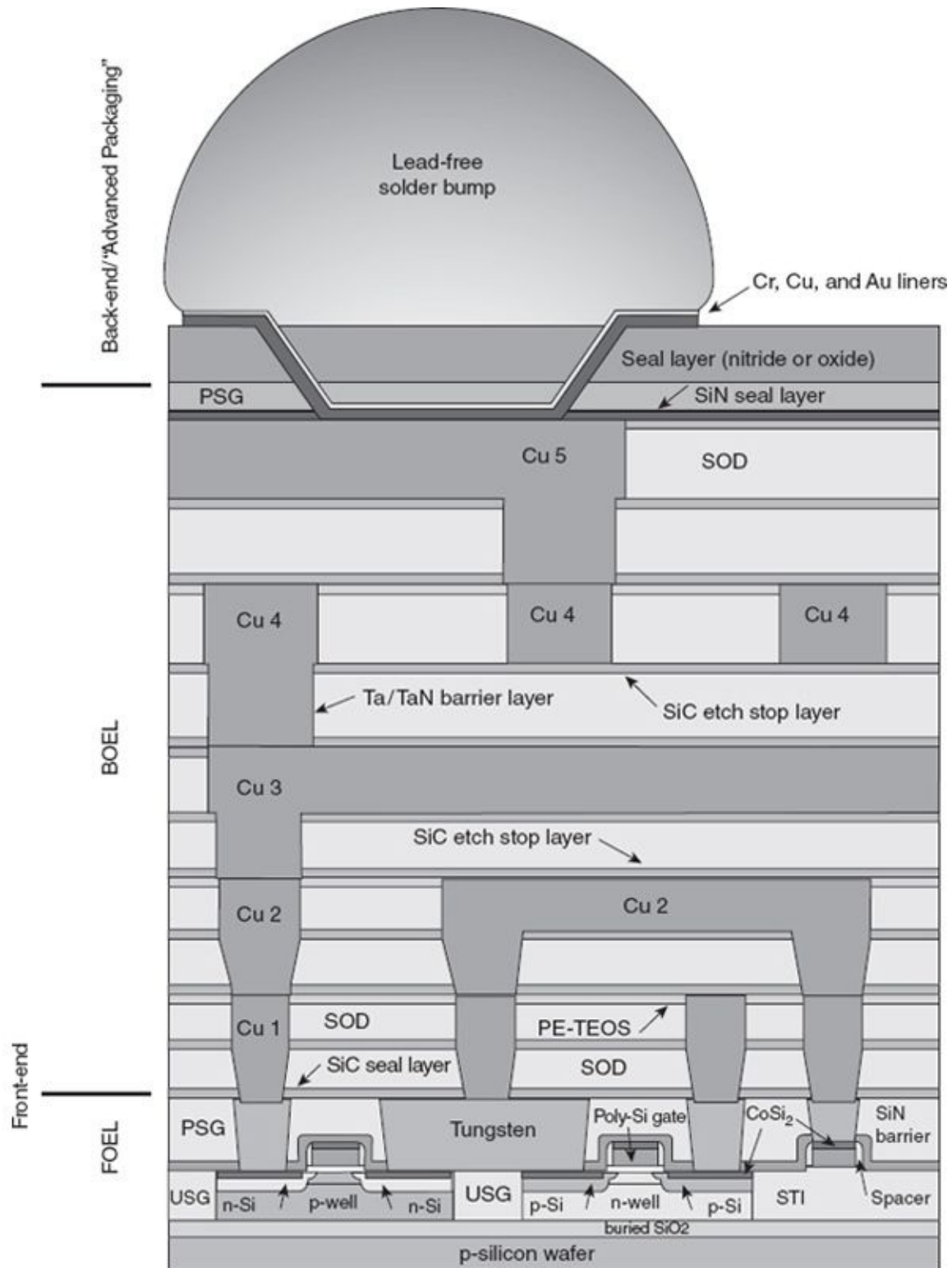


FIGURE 17.13 Five-metal layer cross-section.

But packaging is also ([Chap. 18](#)) exploiting the z plane with three-dimensional

multichip techniques. These can be multiple chips connected together in a package or packages stacked and connected.

Review Topics

Upon completion of this chapter, you should be able to: 1. Explain the concept of binary numbering.

2. List the three major integrated circuit functions.
3. Compare the basis of analog and digital logic circuits.
4. List the user and production advantages of logic gate arrays and PROM circuits.
5. Explain the two major memory circuit types.
6. Make a list of the four nonvolatile memory circuits.
7. Compare the performance and cost factors of DRAM and SRAM memory circuits.

References

1. McConnell, M., "An Experimental 4-Mb Flash EEPROM with Sector Erase," *IEEE Journal of Solid State Circuits*. 26(4), Apr. 1991.
2. Jones, R., "Integration of Ferroelectric Nonvolatile Memories," *Solid State Technology*, Oct. 1997:201.
3. Hu, C., "MOSFET Scaling in the Next Decade and Beyond," *Semiconductor International*, Jun. 1994:105.

CHAPTER 18

Packaging

Introduction

After wafer sort, the chips are still part of the wafer. For use in a circuit or electronic product, they must be separated from the wafer and, in most cases, put in a protective package ([Fig 18.1](#)). As chip component density has grown so has the sophistication of their packages and package processing. Individual packages are typically “cans” for discrete devices and in-line packages for individual ICs. But the exploding mobile wireless devices has required multicircuit functions on a single chip (System on Chip) and also stacking individual chips in the same package in a three-dimensional arrangement. Some of these packaging schemes do away with the traditional hard-shell approach of cans and in-line enclosures.

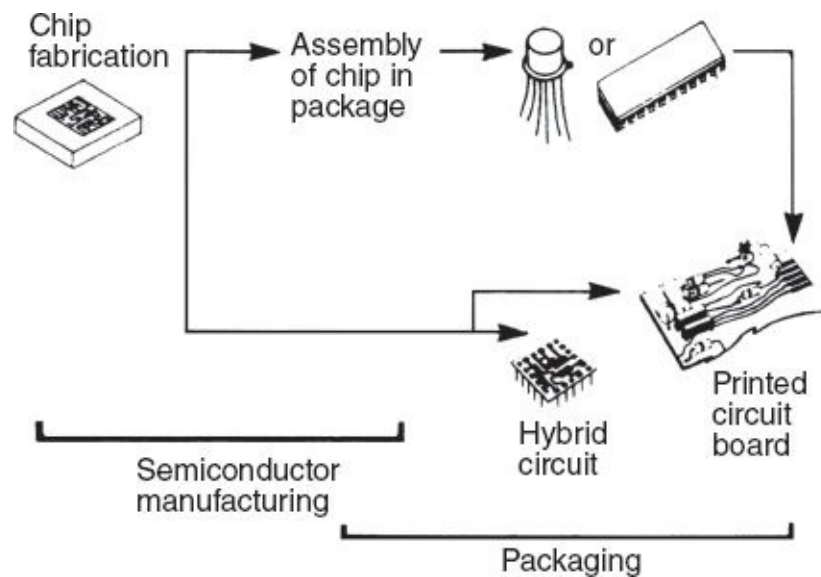


FIGURE 18.1 Chip or packaging process flow diagram.

Chips may also be mounted onto the surface of a ceramic substrate as part of a hybrid circuit, put into a large package with other chips as part of a multichip module (MCM), or be connected directly onboard a printed circuit, chip-on-board, or direct chip attach (COB or DCA) ([Fig. 18.1](#)). All these options share some common processes. The package, in addition to protecting the chip, provides an electrical connection system allowing the chip to be integrated into a larger electronic system, and it provides environmental protection and heat dissipation. This series of processes is known variously as *packaging*, *assembly*, or the *back-end process*. In the packaging process, the chips are called *dies* or *dice*.

Over the years, semiconductor packaging has lagged wafer fabrication in process sophistication and manufacturing demands. The advent of the VLSI/ULSI era in chip density has forced a radical upgrading of chip packaging technology and

production automation. Higher-density chips require many more input connections (I) and many more output connections (O). These are referred to as the I/O count or simply the pin count. The 2007 *Technology Roadmap* (ITRS) projects pin counts increased up to the 4000 to 8000 range in the 2015 time frame ([Fig. 18.2](#)). The ITRS lists pin count, cost, chip size, thickness, and temperature considerations as the primary physical drivers of packaging technology. A packaged (or connected) IC is an electrical system. The thrust is to create increased electrical functioning in chip or package systems. They fall under the general term *System in a Package* (SIP). There are different strategies for SIP systems. Higher package pin counts, like ICs, require spacing the individual leads closer together (called pitch). While the dual in-line package is still the industry's most-used package, demands for newer product are driving innovations in chip and packaging design. Higher pin counts have led to the adoption of bump or flip chip technology. Size and speed considerations have driven the use of chip-scale packages in consumer products, such as cell phones and handheld products. Increasing per package capacity and speed have led to a series of schemes under the name of 3-D packaging. The harsh environments of space, automotive use, and military applications require special packages, processing, and testing to ensure high reliability. These packages, processes, and tests are referred to as *hi-rel*. The other chips and packages are referred to as *commercial* parts.

Package pin count maximum	2012	2013	2014	2015
Low cost	188–1000	198–1050	207–1100	218–1150
Cost performance	720–3367	800–3704	800–4075	880–4482
High performance (FPGA)	5348	5616	5896	6191

FIGURE 18.2 Pin-count predictions. (2007 ITRS.)

No longer is packaging the stepchild of the semiconductor industry. Many feel that, eventually, packaging will be the limiting factor on the growth of chip size. For the time being, however, much effort is going into new package designs, new material development, and faster and more reliable packaging processes.

Chip Characteristics

Throughout this text, many characteristics of discrete devices and integrated circuits have been mentioned. Several of them have a direct bearing on package design and the packaging processes ([Fig. 18.3](#)). The chip density (integration level) determines the number of connections required, with higher-density chips having larger surface areas and more bonding pads. The trend to larger chips has resulted in the need for thicker and larger-diameter wafers. These factors have caused changes in die-separation processes, package design, and the need for wafer thinning.

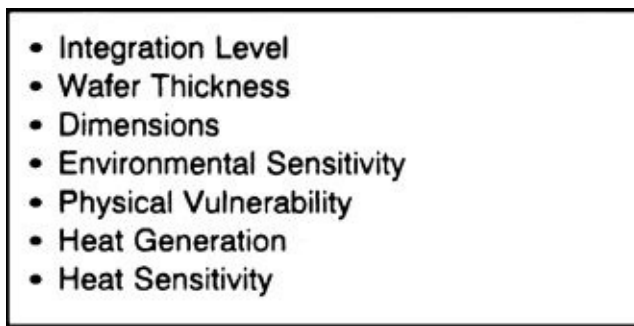
- 
- Integration Level
 - Wafer Thickness
 - Dimensions
 - Environmental Sensitivity
 - Physical Vulnerability
 - Heat Generation
 - Heat Sensitivity

FIGURE 18.3 Chip characterizations affecting the packaging process.

In previous chapters, it was pointed out that the functioning of the chip components (transistors, diodes, capacitors, resistors, and fuses) can be altered by various contaminants. Primary among them are chemicals such as sodium and chlorine. Additionally, other chemicals can attack the chip layers, including environmental factors. For example, particulates, humidity, and static can ruin chips or change their performance. Other concerns are the influence of light and radiation impinging on the chip surface. Some chips are extremely light or radiation sensitive. These factors are considered in the selection of package materials and processing. A dominant chip characteristic is the extreme vulnerability of its surface to physical abuse. The surface components are only a small distance down into the wafer surface, and the surface wiring is thin and vulnerable.

These environmental and physical concerns are addressed in two ways. First is the passivation layer deposited near the end of the fabrication process. This may be a hard layer based on silicon dioxide or silicon nitride. Often, passivation layers are doped with boron, phosphorus, or both to increase their protective properties. Alternatively, it may be a soft layer such as a polyimide ([Fig. 18.4](#)). The second method of protecting the chip is provided by a package itself.

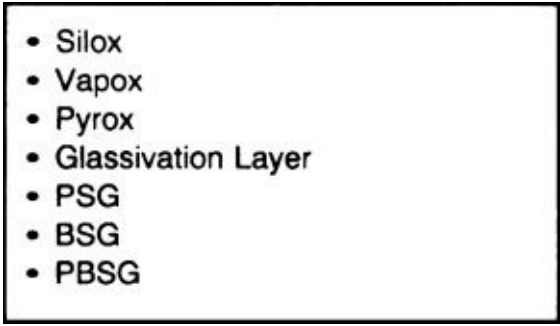
- 
- Silox
 - Vapox
 - Pyrox
 - Glassivation Layer
 - PSG
 - BSG
 - PBSG
-

FIGURE 18.4 Passivation layer types.

Another chip characteristic of importance to the package design and material is heat generation. Chips used in high-power circuits and highly integrated circuits can generate enough heat to actually damage the circuit. Package design includes heat-dissipation factors. Heat is also an important parameter in packaging processes. With ICs using aluminum leads, packaging process temperatures are limited to 450°C to prevent the aluminum and underlying silicon forming an alloy.

Package Functions and Design

There are four basic functions performed by a semiconductor package. They are to provide:

1. A substantial lead system
2. Physical protection
3. Environmental protection
4. Heat dissipation

Substantial Lead System

The primary function of the package is to allow connection of the chip to a circuit board or directly to an electronic product. This connection cannot be made directly, from the thin and fragile on-chip metal system. The metal leads are typically less than 1.5- μm thick and only 1- μm wide. The thinnest bonding wires available are 0.7 to 1.0 mils in diameter, which is many times larger than the on-chip surface wiring. Solder balls used in bump connection technology are about 100 μm in diameter. This difference in scale between the on-wafer wiring sizes and interconnection dimensions is the reason for bonding pads.

Even though the wires are larger, at 1 mil in diameter, they too are very fragile. This fragility is overcome by connecting the chip bonding pads to the outside world with traditional leads, pins, or the balls used in grid array packages ([Fig. 18.5](#)). The lead system is an integral part of the package.

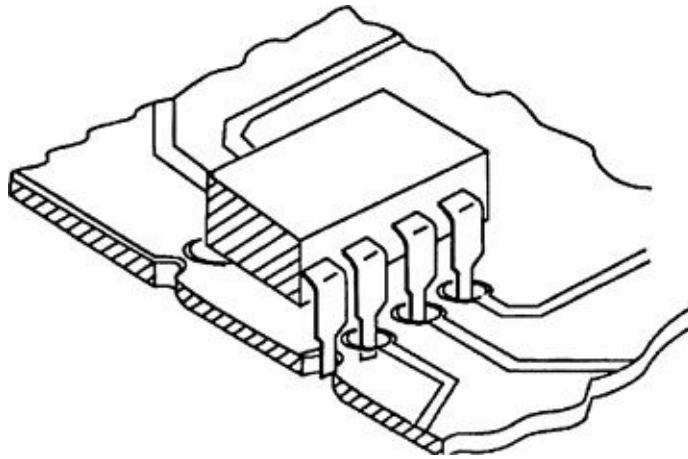


FIGURE 18.5 Dual in-line package (DIP) through-hole assembly.

Physical Protection

The second function of the package is the physical protection of the chip from breakage, particulate contamination, and abuse. Physical protection needs vary from low, as in the case of consumer products, to very stringent, as in the case of automobile circuits, space rockets, and military uses. The protection function is accomplished by securing the chip to a die-attachment area and surrounding the chip, and inner package leads with an appropriate enclosure. The size and eventual use of the chip dictate the choice of materials for the enclosure and the design and size of the package.

Environmental Protection

Environmental protection of the chip from chemicals, moisture, and gases that may interfere with the chip functioning is provided by the package enclosure. It drives the use of “clean” materials and contamination-free processing areas.

Heat Dissipation

Every semiconductor chip generates some heat during operation. Some generate large quantities. The package enclosure materials serve to draw the heat away from most chips. Indeed, one of the factors in choosing a packaging material is its thermal dissipation property. The chips that generate large quantities of heat require additional consideration in the package design. This consideration will influence the size of the package and will often require the addition of metal heat-dissipating fins or blocks on the package.

Common Package Parts

The four functions of a chip package are accomplished through the use of a wide variety of package designs. However, most packages have five common parts. They are as described below using a dual in-line package as an example.

Die-Attachment Area

In the center of the package is an area where the chip is securely attached into the package. This *die-attachment area* may have an electrical connection that serves to connect the back of the chip to the rest of the lead system. A major requirement for this area is absolute flatness so as to intimately support the chip in the package without cracking or bending ([Fig. 18.6](#)).

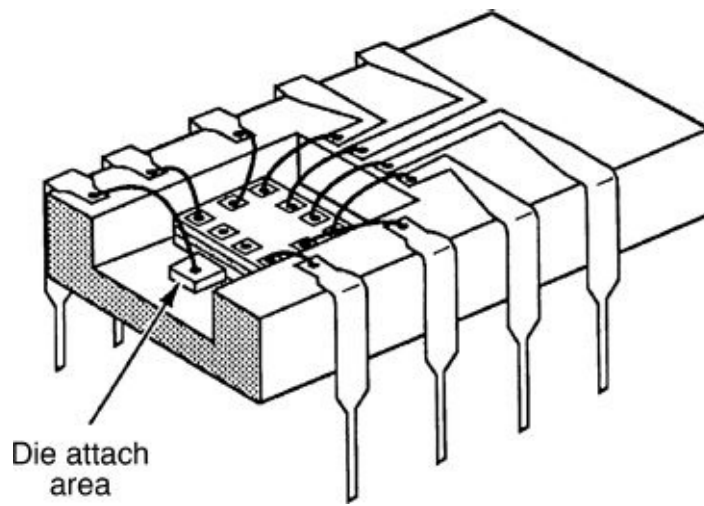


FIGURE 18.6 Die attach area of package.

Inner and Outer Leads

The metal lead system chip is continuous from the die-attach cavity to the printed circuit board (PCB) or electronic product. The inner connections connect the chip to the package and are called the *inner leads*, *bonding lead tips*, or *bond fingers*. The inner leads are generally the narrowest portion of the package lead system. The *outer leads* connect the package to the outside world ([Fig. 18.7](#)). Most of the lead systems are composed of one continuous piece of metal. One exception is the side-brazed package. In this package construction method, the outer leads are brazed onto the interior leads. Two different alloys are used for the outer lead system—either an iron-nickel alloy or a copper alloy. The iron-nickel alloy is desirable for its strength and stability, while copper is used for its electrical and heat-conduction properties. Tape automated bonding (TAB) schemes have one lead, per bonding pad, from the chip that is in turn bonded directly to a circuit board. Flip chip packages use solder bumps or balls for the inner connections ([Fig. 18.8](#)).

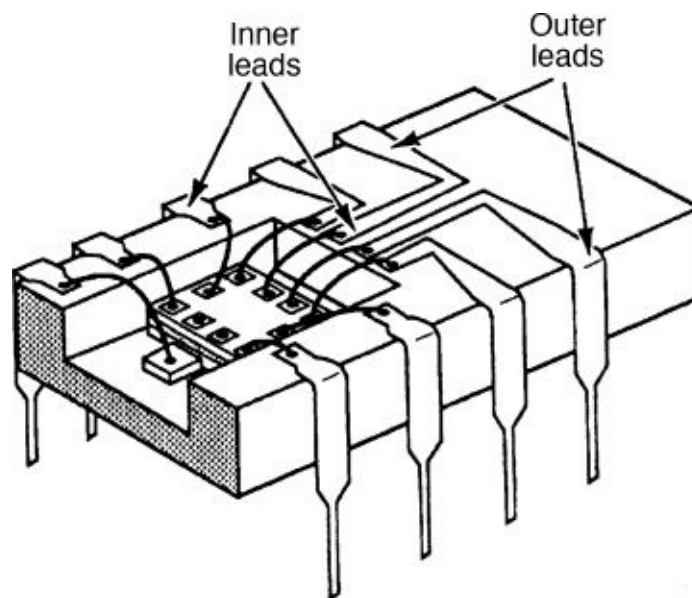


FIGURE 18.7 Inner and outer leads.

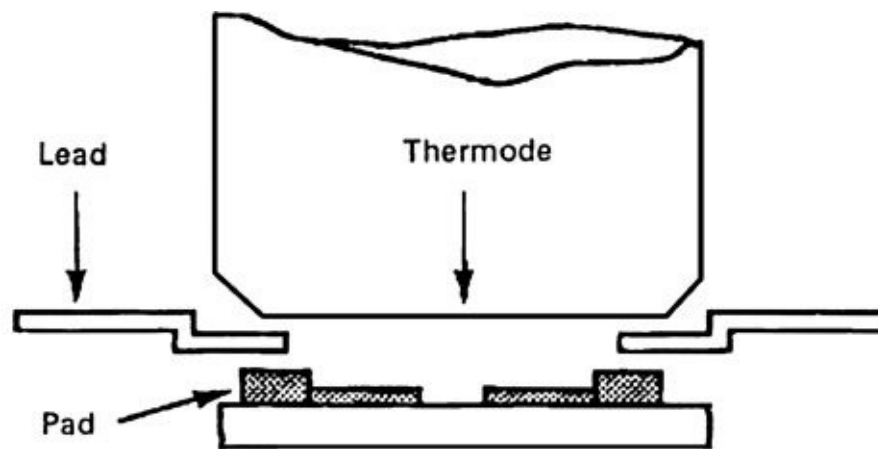


FIGURE 18.8 Tape Automated Bonding. (Source: Microelectronics Handbook, by Tummala and Rymaszewski.)

Enclosures

The die-attach area, bonding wires, and inner and outer leads constitute the electrical parts of the package. The other part is the *enclosure* or *body*. This is the part that provides the protection and heat-dissipation functions. These functions are achieved by several different techniques and package designs as described in the sealing section. Seals fall into two categories: hermetic and nonhermetic ([Fig. 18.9](#)).

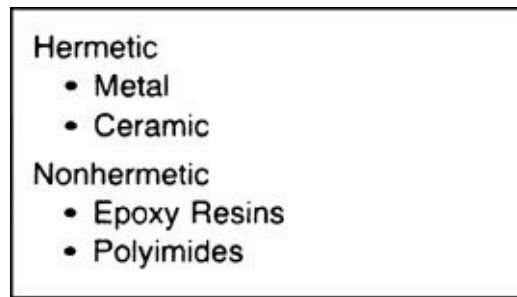


FIGURE 18.9 Package sealing designations.

Hermetic sealing results in a package that is impervious to the penetration of moisture and other gases. Hermetic seals are required for chips operating in harsh and demanding environments such as in rockets and space satellites. Metal and ceramic enclosures are the preferred materials for hermetically sealed packages.

Nonhermetically sealed packages are adequate for most consumer applications such as computers and entertainment systems. This sealing system provides good and adequate environmental protection of the chip, in normal operating situations. A better term for this type of enclosure sealing method would be *less hermetic*. Nonhermetic packages are composed of epoxy resins or polyimide materials and are generally referred to as *plastic packages*.

Cleanliness and Static Control

The chips are vulnerable to contamination during their entire lifetime. Assembly areas are not generally required to maintain the same cleanliness levels as fabrication areas (see [Chap. 4](#)). High-rel areas, particularly, demand higher cleanliness levels. In fact, many companies are finding that halfway contamination control programs are doomed to failure. Consequently, more assembly areas are practicing very stringent controls, especially for people-generated particles and chemicals.

An environmental danger that is most serious in packaging areas is static. In the fabrication cleanrooms, static is controlled primarily to prevent the attraction of particles to the wafer surface. This is also a concern in a packaging area. But the greatest concern is *electrostatic discharge (ESD)*. Static charge can build up to levels of tens of thousands of volts. If this voltage is suddenly discharged onto a chip surface, it can easily destroy a portion of the circuit. MOS gate structures are particularly vulnerable to ESD. [Figure 18.10](#) lists common static-control practices used in packaging areas.

- 
- HEPA Filters/VLF Air
 - Smocks, Hats, Shoe Covers
 - Finger Cots or Gloves
 - Filtered Chemicals
 - Sticky Floor Mats
 - Static Control

FIGURE 18.10 Contamination control practices.

Every packaging area making high-density chips will have an active antistatic program ([Fig. 18.11](#)). The antistatic program includes operators wearing grounding straps and nonstatic smocks; the use of antistatic carrier materials; moving work by lifting rather than sliding; and grounded equipment, work surfaces, and floor mats. Static is also reduced by the placement of ionizers in nitrogen and air blow-off guns ([Fig. 18.12](#)) in the path of air coming out of High-efficiency particulate air (HEPA) filters.

- Wrist Ground Straps
- Nonstatic Garments
- Antistatic Materials
- Grounded Equipment
- Grounded Work Surfaces
- Grounded Floors or Floor Mats

FIGURE 18.11 Static control practices.

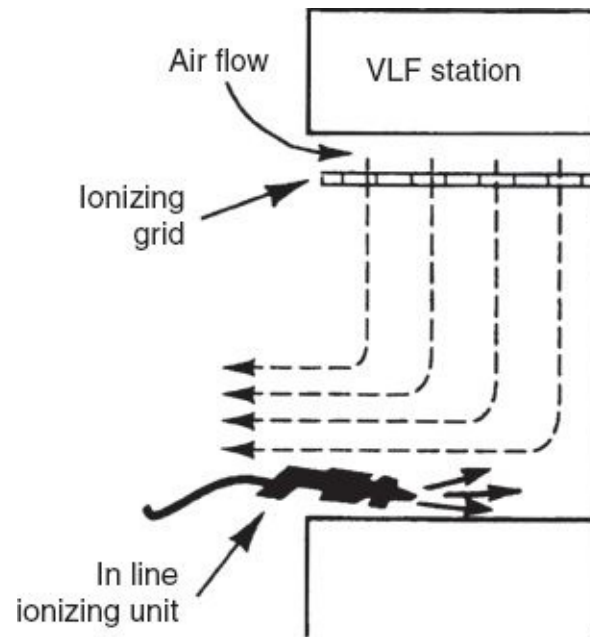


FIGURE 18.12 Static control techniques.

Basic Bonding Processes

In wafer fabrication, the wafers pass in and out of four basic operations (layering, patterning, doping, and heat processing) numerous times. In packaging, as well, there are also several basic operations ([Fig. 18.13](#)). However, packaging is generally a once-through process. Each of the major processes is required only once. As in fabrication, the exact order of the operations is determined by the product package type and other factors. Each process may or may not be used in a particular process, and each is customized to the particular chip and package requirements. There are three standard methods of connecting (bonding) the chip to the package: wire bonding, tape automated bonding (TAB) and bump (ball) or flip chip technology ([Fig. 18.14](#)). Wire bonding is a well-established process used for most single chips in an individual package. It also finds use in some chip direct on a circuit board and in connection chips in three-dimensional packages. TAB is a process of connecting conductive leads onto a lead frame that in turn is soldered directly on a circuit board. Bump or ball bonding requires that solder balls are positioned over the bonding pads on a chip, but while the wafer is still in the wafer-fabrication process area. One application is *wafer scale packaging* (WSP). The actual bonding to the package requires the wafer to be flipped upside down and heat processed to make the electrical connections. This process is called *bump*, *ball*, and/or *flipchip* technology. Three-dimensional packaging uses variations on the bump, ball, or flipchip schemes. This section will describe common package characteristics followed by three basic bonding processes. The different packages used are described in the Sealing Techniques section.

• Backside Preparation	• Preseal Inspection
• Die Separation	• Packaging Sealing
• Die Pick	• Plating
• Die Attach	• Lead Trim
• Inspection	• Marking
• Bonding	• Final Tests

FIGURE 18.13 Basic packaging operations.

- Wire Bonding
 - Gold
 - Aluminum
- Tape Automated Bonding (TAB)
- Bump/Ball - Flip Chip

FIGURE 18.14 Overview of bonding techniques.

Wire Bonding Process

Prebonding Wafer Preparation

After the final passivation layer and an alloy step in wafer fabrication, the circuits are complete. However, one or two additional processes may be performed on the wafer before transfer to packaging. These steps (wafer thinning and backside gold) are optional, depending on the wafer thickness and the particular circuit design.

Wafer Thinning

Wafer scale processes use the bump or ball bonding process. The balls are soldered onto the chip bonding pads at the end of the usual wafer-fabrication process. This process is described in the bonding sections.

The trend to thicker wafers presents several problems in the packaging process. Thicker wafers require a more expensive complete saw-through method at die separation. While sawing produces a higher-quality die edge, the process is more expensive in time and consumption of diamond-tipped saws. Thicker die also require deeper die-attach cavities, resulting in a more expensive package. Both of these undesirable results are avoided by thinning the wafers before die separation.

Another situation requiring wafer thinning is electrical in nature. If the wafer backs are not protected as the wafers go through the dopant operations in fabrication, the dopants will form electrical junctions in the wafer back, which may interfere with good conduction in the back contact that is required for the circuit to operate correctly. These junctions may require physical removal by wafer thinning. Additionally wafer scale and three-dimensional packaging requires thinner wafers to avoid thicker packages and to increase electrical functioning.

The thinning step generally takes place between wafer sort and die separation. Wafers are reduced to a thickness in the 100- μm range.¹ It is done by the same processes (mechanical grinding and chemical-mechanical polishing or CMP) used to grind wafers in the wafer-preparation stage. Flipchip packaging also requires wafer thinning.

Wafer thinning is a worrisome process. In back grinding, there is the concern of scratching the front of the wafer and of wafer breakage. Since the wafer must be held down on the grinder or polishing surface, the front of the wafer must be protected and, once thinned, wafers are easier to break. In back etching, there is a similar need to protect the front of the wafer from the etchant. The protection can be provided by spinning a thick layer of photoresist on the front side. Other methods include attaching adhesive-backed polymer sheets cut to the wafer diameter. Stresses induced in the wafer by grinding or polishing processes must be controlled to prevent wafer and die warping. Wafer warping interferes with the die-separation process (broken and cracked die). Die warping creates die attach problems in the packaging process.²

Backside Gold

Another optional wafer process is adding a layer of backside gold. A layer of gold is required on wafers that are going to be attached to the package by eutectic techniques (see the “Die Attach” section). The gold is usually applied in the fabrication area (after back grinding) by sputtering.

Die Separation

The traditional chip-packaging process starts with the separation of the wafer into individual dies. The two methods of die separation are scribing and sawing ([Fig. 18.15](#)). For bump or ball processes, the die separation occurs after the bump or ball system is established on the wafer.

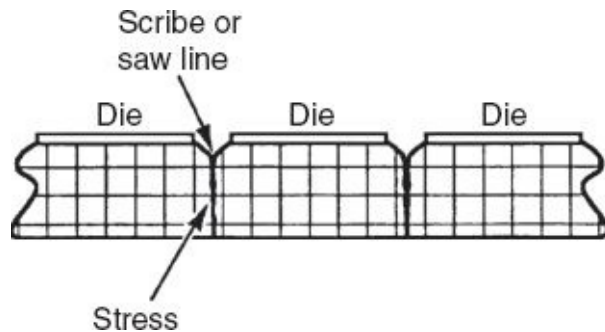


FIGURE 18.15 Scribe and saw separation.

Sawing

Thicker wafers have led to the development of sawing as the preferred die-separation method. A saw consists of a wafer table with rotation capability, a manual or automatic vision system for orienting the scribe lines, and a diamond-impregnated round saw. Two techniques are used. Both start with the passage of the diamond saw over the scribe lines. For thinner wafers, the saw is lowered into the wafer surface to create a trench about one-third of the way through the wafer. The separation of the wafer into die is completed by the stress and roller technique used in the scribing method. The second sawing method is to separate the die by a complete saw-through of the wafer.

Often, the wafers for complete saw-through are first mounted on a flexible plastic film. The film holds the die in place after the sawing operation and aids the die pick operation. Sawing is the preferred die separation method due to the cleaner die edges and the fewer cracks and chips left on the sides of the die ([Fig. 18.16](#)).

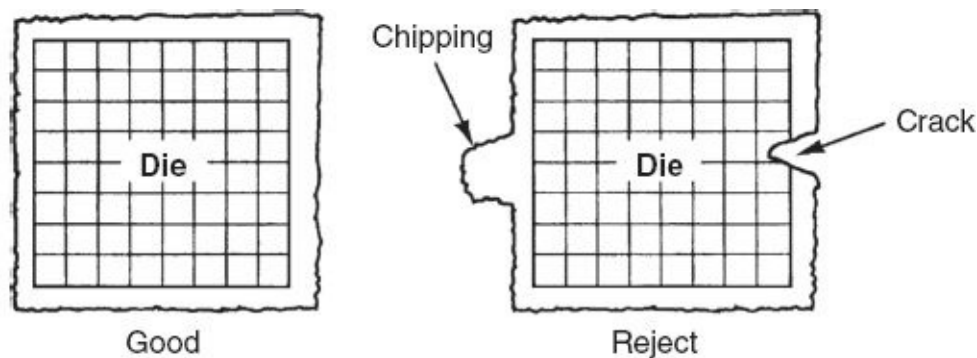


FIGURE 18.16 Die-separation results.

Scribing

Scribing, or *diamond scribing*, was the first production die-separation technique developed in the industry. It requires dragging a diamond-tipped scribe through the center of the scribe lines and separating the die by flexing the wafer. Scribing becomes unreliable in wafers over 10 mils thick.

Die Pick and Place

After sawing, the separated die are transferred to a station to select (pick) the functioning die. In operation, a memory tape or disk that has recorded the locations of the good die (from wafer sort) and is loaded into an automatic “pick” tool. A vacuum wand picks up the good die and automatically places them in a sectioned plate for transfer to the next operation ([Fig. 18.17](#)).

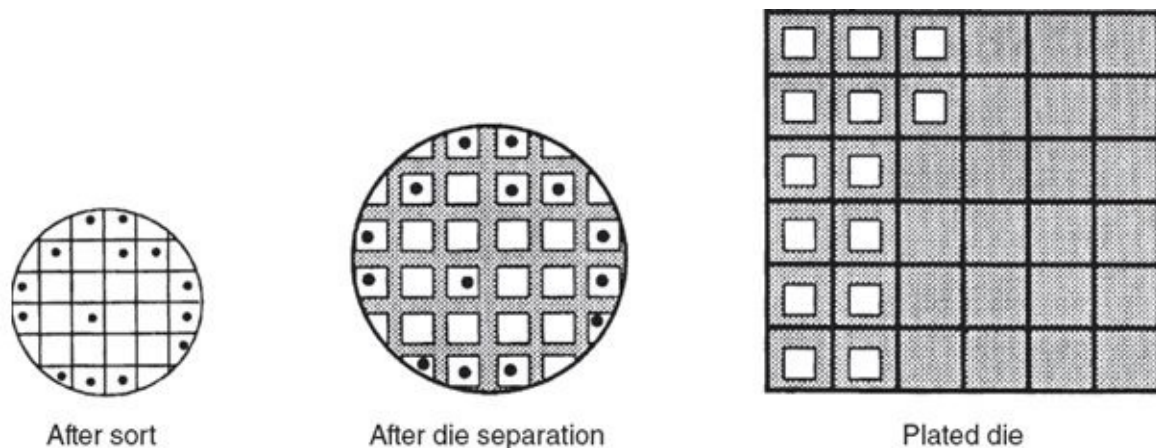


FIGURE 18.17 Die separation to plate.

In the manual method, an operator will pick up each of the non-inked dies with a vacuum wand and place it in a sectioned plate. Wafers that come to the station on the flexible film are first placed on a frame that stretches the film. The stretching separates the die, which aids the die pick operation.

Die Inspection

Before being committed to the rest of the process, the die are given an optical inspection. Of primary interest is the quality of the die edge, which should be free of chips and cracks. This inspection also sorts out surface irregularities, such as scratches and contamination. The identification of damaged die saves the expense and time of packaging a failed die.

Die Attach

Die attachment has several goals: to create a strong physical bond between the chip and the package, to provide either an electrical conducting or insulating contact between the die and the package, and to serve as a medium to transfer heat from the chip to the package. Flipchip schemes do not have a die-attach step.

The die-attach bond should not loosen or deteriorate during subsequent processing steps or when the package is in use in an electronic product. This requirement is especially important for packages that will be subjected to high physical forces, such as in rockets. Additionally, the die-attach materials should be contaminant-free and should not outgas during subsequent process heating steps. Lastly, the process itself should be productive and economical.

Eutectic Die Attach

There are two principal methods of die attach: *eutectic* and *epoxy* adhesives. The eutectic method is named for the phenomenon that takes place when two materials melt together (alloy) at a much lower temperature than either of them separately. For die attach, the two eutectic materials are gold and silicon ([Fig. 18.18](#)). Solder is also used. Gold melts at 1063°C, while silicon melts at 1415°C. When the two are mixed together, they start alloying at about 380°C. Gold is plated onto the die-attach area and alloys with the bottom of the silicon die when heated.



FIGURE 18.18 Die-attach material matrix.

The gold for the die-attach layer is actually a sandwich. The bottom of the die-attach area is deposited or plated with a layer of gold. Sometimes, a preformed piece of metal composed of a gold and silicon mixture is placed in the die-attach area. When heated, these two layers form a thin alloy layer between the wafer back and the package.

Eutectic die attach requires four actions. First is the heating of the package until the gold-silicon forms a liquid. Second is the placement of the chip on the die-attach area. Third is an abrasive action, called *scrubbing*, that forces the die and package surfaces together. It is this action, in the presence of the heat, that forms the gold-silicon eutectic layer. The fourth and final action is the cooling of the system, which completes the physical and electrical attachment of the chip and package.

Eutectic die attach can be performed manually or by an automated machine that performs the four actions. Gold-silicon eutectic die attach is favored for high-reliability devices and circuits for its strong bonds, heat-dissipation properties, thermal stability, and lack of impurities.

Epoxy Die Attach

The alternate die-attach process uses thick liquid epoxy adhesives. These adhesives can form an insulating barrier between the chip and package or be electrically and heat conductive with the addition of metals such as gold or silver. Polyimide may also be used as an adhesive. Popular adhesives are silver-filled epoxy for copper-lead frames and silver-filled polyimide for Alloy 22 metal frames.³

The epoxy process starts with the deposit of the epoxy adhesive in the die-attach area by dispensing the adhesive from a needle or screen printing it into the die-attach area. The die, held by a vacuum wand, is positioned in the center of the die-attach area. The second action is to force the die into the epoxy to form a thin uniform layer under the die. The final action is a curing step in an oven at an elevated temperature that sets the epoxy bond.

Epoxy die attach is favored for its economy and ease of processing, in that the package does not have to be heated on a stage. This factor makes the automation of the process easier. When compared to gold-silicon eutectic die attach, epoxy has the disadvantage of potential decomposition at the high temperatures of bonding and sealing operations. Epoxy die-attach films also do not have the bonding power of the eutectics.

Regardless of the attachment method used, there are several marks of a successful die attach. One is the proper and consistent alignment of the die in the die-attach area. Proper placement is required for faster and higher-yield automatic bonding. Another desired result is a solid, uniform, and void-free contact over the entire area of the chip. This is necessary for mechanical strength and good thermal conduction. One evidence of a uniform bond is a continuous joint or “fillet” between the die edge and the package. The final mark of a good die-attach process is a die-attach area free of flakes or lumps that can come loose during use and cause a malfunction.

Wire Bonding

Once the die and package are attached, they go to the wire-bonding process. This is perhaps the most critical of all the assembly operations. In wire bonding, up to hundreds of wires must be perfectly bonded from the bonding pads to the package inner leads (Fig. 18.19). The wire bonding procedure is simple in concept. A thin (0.7 to 1.0 mil) wire is bonded to the chip-bonding pad and spanned to the inner lead where a second bonding operation takes place. Last, the wire is clipped and the entire process is repeated at the next bonding pad. While simple in concept and procedure, wire bonding is critical because of the precise wire placement and electrical contact requirements. In addition to accurate placement, each and every wire must make good electrical contact at both ends, span between the pad and inner lead in a prescribed loop without kinks, and be at a safe distance from neighboring wires. Wire loops in conventional packages are 8 to 12 mils, while those in ultra-thin packages are 4 to 5 mils.⁴ Distances between adjacent wires are referred to as the *pitch* of the bonding.

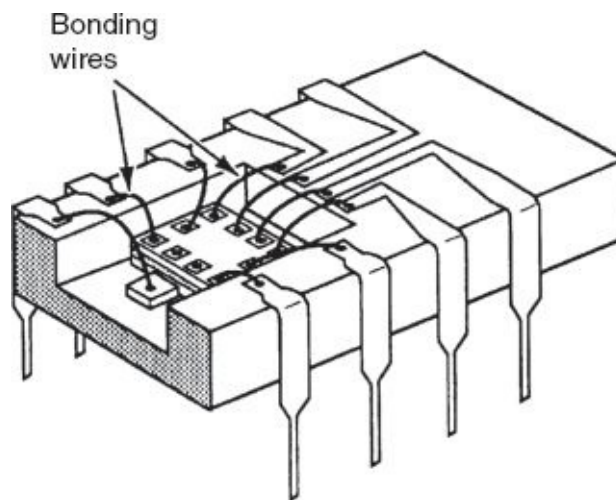


FIGURE 18.19 Bonding wire.

Wire bonding is done with either gold or aluminum wires. Both are highly conductive and ductile enough to withstand deformation during the bonding steps and still remain strong and reliable. Each has its advantages and disadvantages, and each is bonded by different methods.

Gold Wire Bonding

Gold has several pluses as a bonding wire material. It is the best known room-temperature conductor and is an excellent heat conductor. It is resistant to oxidation and corrosion, which translates into an ability to be melted to form a strong bond with the aluminum bonding pads without oxidizing during the process. Two methods are used for gold bonding. They are *thermocompression* and *thermosonic*.

Thermocompression bonding (also known as *TC bonding*) starts with the positioning of the package on the bonding chuck and the heating of the chip and package to between 300 and 350°C. Chips that are going to be enclosed in an epoxy molded package are processed through die attach and bonded with the chip on the lead frame only. The bonding wire is fed out of a thin tube called a *capillary* ([Fig. 18.20](#)). An instantaneous electrical spark or small hydrogen flame melts the tip of the wire into a ball and positions the wire over the first bonding pad. The capillary moves downward, forcing the melted ball onto the center of the bonding pad. The effect of the heat (thermal) and the downward pressure (compression) forms a strong alloy bond between the two materials. This type of bonding is often called *ball bonding* (not to be confused with bump or ball bonding). After the ball bond is complete, the capillary feeds out more wire as it travels to the inner lead. At the inner lead, the capillary again travels downward to where the gold wire is forced by the heat and pressure to melt onto the gold-plated inner lead. The spark or flame then severs the wire, forming the ball for the next pad bond. This procedure is repeated until every pad and its corresponding inner pad are connected.

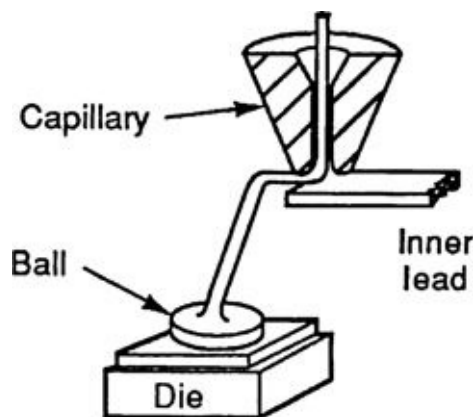


FIGURE 18.20 Gold ball wire bonding.

Thermosonic gold ball bonding follows the same steps as thermocompression bonding. However, it can take place at a lower temperature. This benefit is provided by a pulse of ultrasonic energy that is sent through the capillary into the wire. This

additional energy is sufficient to provide the heat and friction to form a strong alloy bond.

The majority of production gold wire bonding is done on automatic machines that use sophisticated techniques to locate the pads and span the wire to the correct inner lead. The fastest bonding machines can perform thousands of bonds per hour. There are two major drawbacks to the use of gold bonding wires. First is the expense of the gold. Second is an undesirable alloy that can form between the gold and aluminum. This alloy can severely reduce the conduction ability of the bond. It forms a purplish color and is known as the “purple plague.”

Aluminum Wire Bonding

Aluminum wire, while not having the conduction and corrosion-resistance properties of gold, is still an important bonding wire material. A primary advantage is its lower cost. The second advantage is that the bond with the aluminum bonding pad is a monometal system and thus less susceptible to corrosion. Also, aluminum bonding can take place at lower temperatures than gold bonding, making it more compatible with the use of epoxy die-attach adhesives.

The bonding of the aluminum follows the same major steps as gold wire bonding. However, the method of forming the bond is different. No ball is formed. Instead, after the aluminum wire is positioned over the bonding pad, a wedge ([Fig. 18.21](#)) forces the wire onto the pad as a pulse of ultrasonic energy is sent down the wedge to form the bond. After the bond is formed, the wire is spanned to the inner lead where another ultrasonic-assisted wedge bond is formed. This type of bonding is known variously as *ultrasonic* or *wedge bonding*.

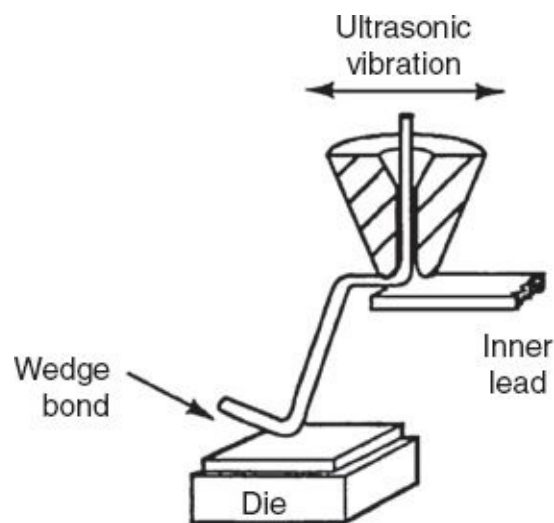


FIGURE 18.21 Aluminum wedge bonding.

After this bond, the wire is cut. At this point in the process, a major difference between the bonding of the two materials occurs. In gold bonding, the capillary moves freely from pad to inner lead, to pad, and so forth, with the package in a fixed position. In aluminum wire bonding, the package must be repositioned for every single bonding step. The repositioning is necessary to line up the pad and inner lead along the direction of travel of the wedge and wire. This requirement places an additional difficulty on the designers of automatic aluminum bonding machines. Nevertheless, most production aluminum bonding is done on high-speed machines.

Tape Automated Bonding Process

Tape automated bonding (TAB) is used to connect chips directly to a PCB when extreme thinness is required, such as in credit-card-size radios. The TAB process starts with formation of the lead system on a flexible strip of tape. Various methods are used to form the lead system. The metal for the system is deposited on the tape by sputtering or evaporation. Formation of the lead system is either by mechanical stamping or patterning techniques similar to the fabrication-patterning process. The result is a continuous tape containing many individual lead systems. For the bonding operation ([Fig. 18.22](#)), the chip is positioned on a chuck, and the tape is moved by sprockets until one of these lead systems is positioned exactly over the chip. In this position, the inner leads of the system should be positioned over the bonding pads of the chip.

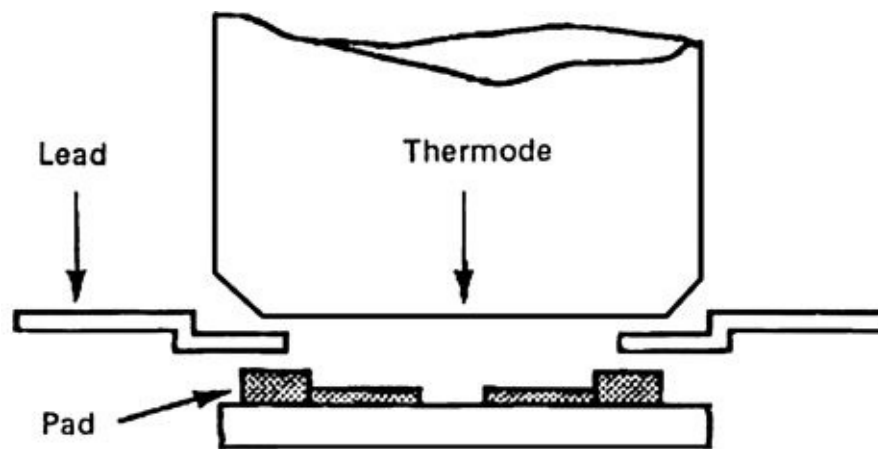


FIGURE 18.22 Tape automated bonding.

The bond is completed with a tool known as a *thermode*. The thermode is faced with a flat diamond surface and is heated. The thermode is moved downward, first contacting the inner leads. It continues downward with enough pressure to force the inner leads onto the chip bonding pads. The heat and the pressure are regulated to cause a physical and electrical bond between the two. Large chips require a larger TAB bonding area. For these chips, the thermode is faced with a synthetic diamond.

TAB bonding is also used to bond package leads to a circuit board. The advantages of TAB are speed, in that all of the bonds to the chip are made in one action, and the ease of automation offered by the tape and sprocket system.

Bump or Ball FlipChip Bonding

Wire bonding presents several problems. There are electrical resistances associated with each bond. There are minimum height limits imposed by the required wire loops. There is the chance of electrical performance problems or shorting if the wires come too close to each other. Plus, the wires require an individual bonding step at both the chip bonding pad and at the package lead. Perhaps the biggest problem results from the increasing number of connections (pin count) needed to operate larger circuits. Chip designers simply run out of space to locate the required number of connections around the periphery of the chip. These issues are addressed by replacing wires with a deposited metal *bump* on each bonding pad. The bumps are also called *balls*, as in naming packages using bump or flipchip processes as ball-grid arrays (BGAs). This bonding method allows chip design with bonding pads located both along the edge of the die as well as in the interior of the die ([Fig. 18.23](#)). These locations place the bump closer to the chip circuitry, increasing signal processing speed. Connection to the package is made when the chip is flipped over and the bump soldered to a corresponding package inner lead ([Fig. 18.24](#)) on a package or PCB. IBM calls their version of this technology controlled collapse chip connection (C4).⁵

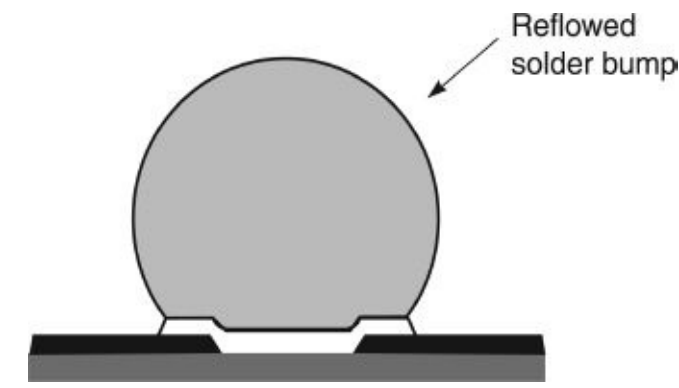


FIGURE 18.23 Reflowed solder bump.

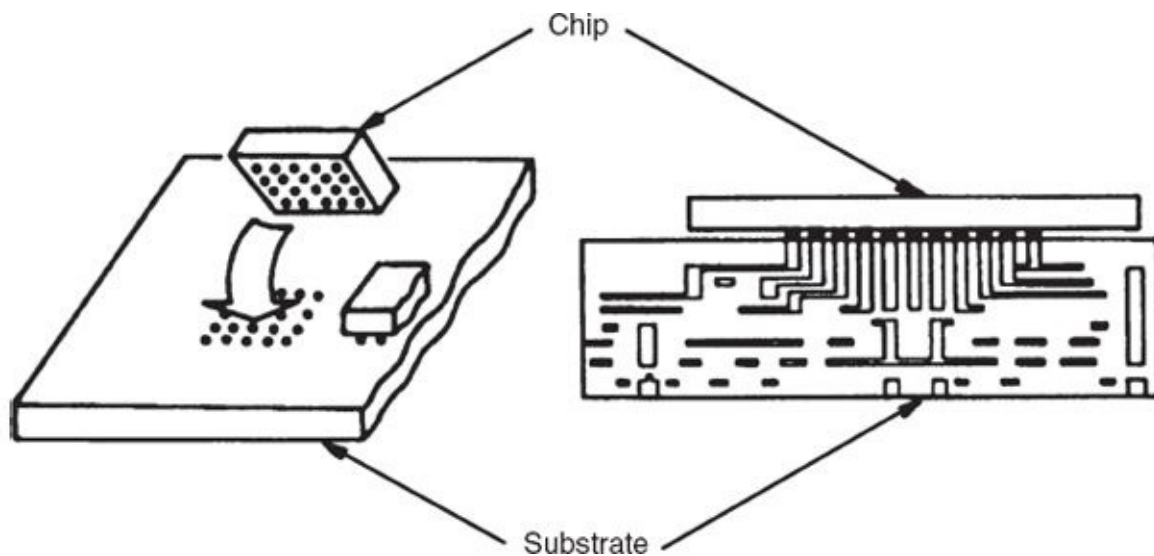


FIGURE 18.24 Flipchip joining. (Source: Microelectronics Handbook, by Tummala and Rymaszewski.)

This process leaves the die suspended above the package surface. Physical stresses and strains are absorbed by the soft solder bump. Additional stress tolerance is provided by filling the gap with an epoxy filling, called an underfill.

Flip chip-connection technology starts in the wafer-fabrication process. Wafers are processed through the usual metallization, passivation, and bonding pad-patterning processes. Thinning is required for chips headed to multichip packages. Wafers for three-dimensional packaging are typically thinned to 75- μm thickness.⁶

A number of process flows are available to form the solder bumps on the bonding pad. The process described below is an example.

Sputter Deposit Intermetal Stack

Lead or tin solder balls are the preferred “bump” material. However, the final metal layers in most ICs is aluminum that has the drawbacks of easily oxidizing to aluminum oxide which is an insulator. Electrically connecting the aluminum and bump ball requires an under-bump-metallization (UBM) stack (also called a plug) of metals.² The stack materials must also bond to the sides of the via hole essentially sealing the underlying IC pad for contamination. The UBM process starts with a removal of the aluminum oxide layer by sputter etch or a wet chemical treatment.

A typical metal stack has four layers:

- An adhesion layer to “wet” the silicon (Ti/Cr/Al)
 - A diffusion barrier to prevent unwanted metallic contamination from diffusing into the IC (Cr:Cu)
 - A solderable layer (Cu/Ni:V)
 - An oxidation barrier layer to provide a mechanical and electrical connection with the bump metal (Au)
-

Example Bump or Ball Process

[Figure 18.25](#) shows an alternative process. It starts with a blanket sputter deposition of a Ti-Ni or copper. After a patterning step over the bonding pad, there is another deposit of Ni, followed plating of the solder plug. The system is completed with the removal of the resist and etching away the intermetal materials left on the surface. The end result is similar to the first process—a ball of lead-tin solder positioned on the bonding pad.

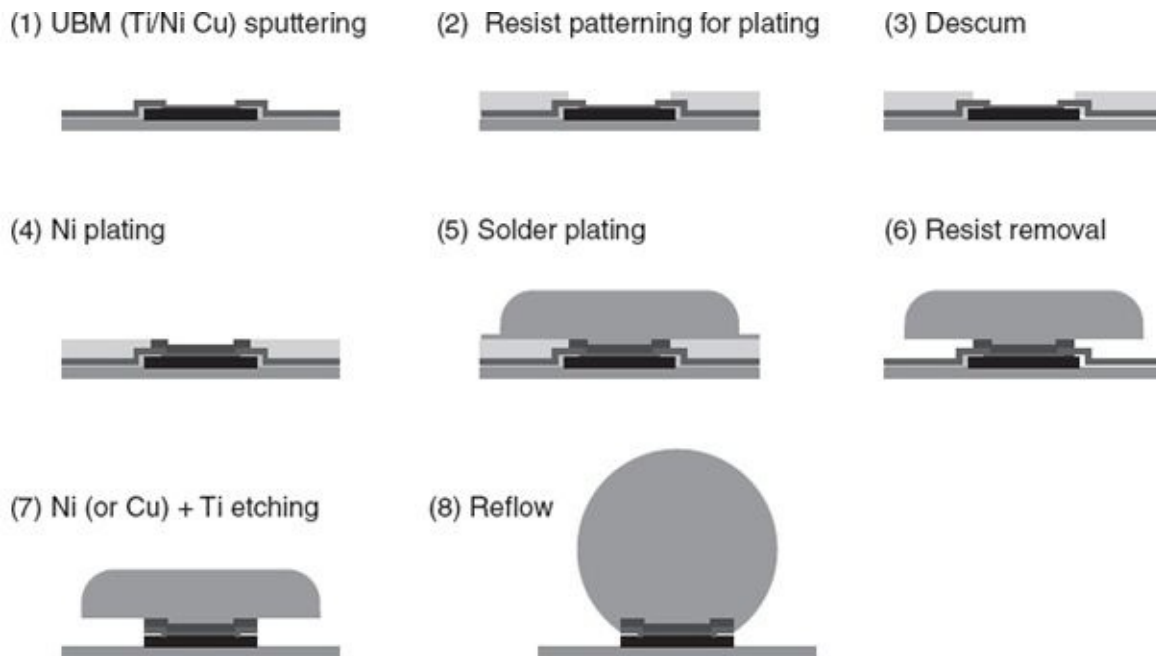


FIGURE 18.25 Alternative bump process. (Courtesy of Future Fab International.)

Copper Metallization (Damascene) Bump Bonding

Copper metallization from the damascene process also requires intermediate layers as shown in [Fig. 18.26](#).

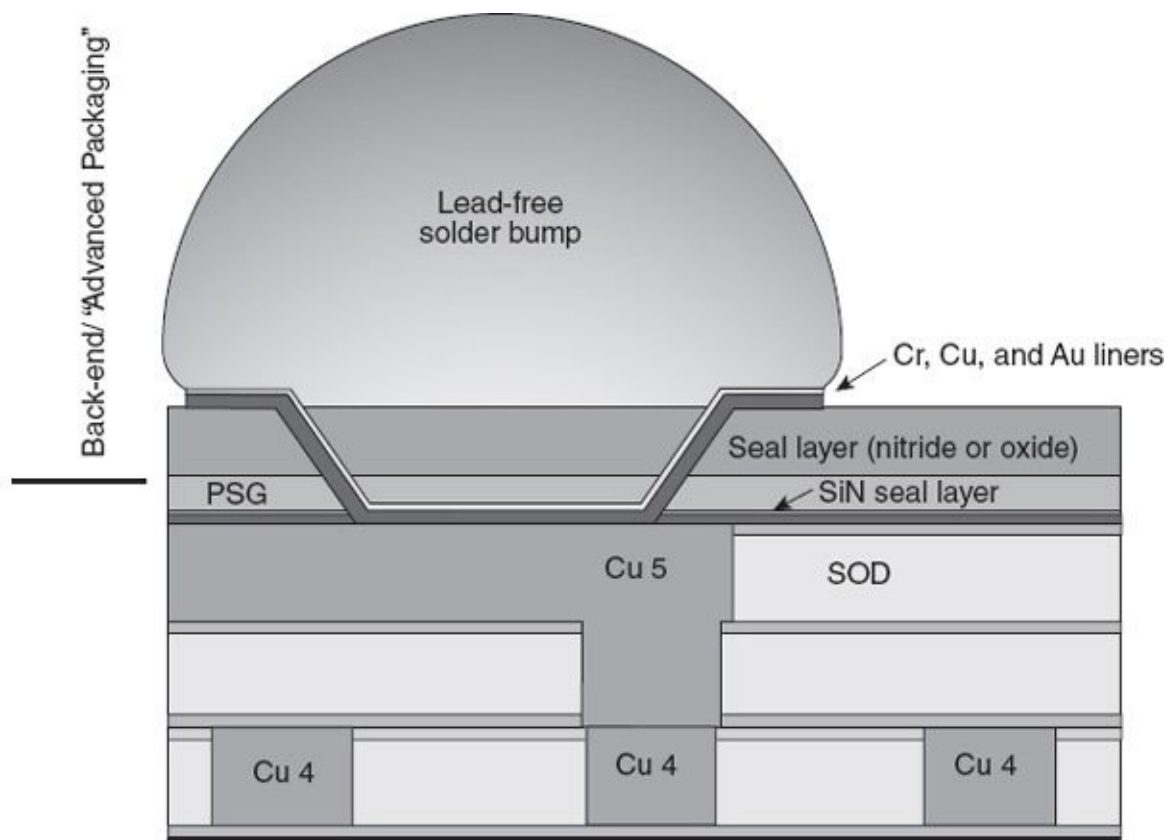


FIGURE 18.26 Bonding ball on copper metallization.

Reflow

A heating of the wafer in the 350°C range in hydrogen enhances the electrical contact to the intermetal stack. Also, surface tension pulls the solder deposit into a spherical shape, much like surface tension causes spherical soap bubbles.

Die Separation and Die Pick and Place

These two steps are the same as in the wire-bonding process. Typically, the die are affixed to a flexible tape on a reel that allows automatic handling in the subsequent processes.

Alignment of Die to Package

Packages used in flipchip systems are usually ceramic or organic. *Organic* is the term used for packages based on the same materials used for circuit boards. Sometimes they are referred to as “plastic” packages. The first step is flipping the die upside down and aligning it to the package such that the bumps on the die are positioned over the corresponding leads on the package.

Attachment to Package (or Substrate)

In the presence of a solder flux, the die-package combination is heated in an oven to melt the solder ball to the package lead.

Deflux

A cleaning process removes excess flux from the surface.

Underfillment

An epoxy is introduced under the corners of the die soldered to the package. When heated, the epoxy is drawn by capillary action underneath the entire chip. After a curing step, it provides a “cushion” to absorb additional stresses on the system. The package designs using the bump or flipchip technique are discussed later on in this chapter.

Encapsulation

Typically, a molding compound is deposited over the upside-down chip to provide environmental and physical protection. The combination of the substrate, die, and protective top layer form the functions of a tradition package.

Postbonding and Preseal Inspection

After bonding, a number of steps are required to complete the packaging operation. The packaging steps described follow the traditional wire bonding and individual package process flow. Most of them have to be performed at some point in the bump or flipchip and tape automated bonding processes. An important step in the traditional chip-packaging process is the preseal or precap inspection, sometimes called *third optical inspection*, which takes place after the bonding step. The inspection is performed to provide feedback on the quality of the operations already performed. It also is performed to reject packaged chips that may represent reliability hazards when the chip is in operation in the field.

While there are many levels of inspection criteria, they fall into two main categories: commercial and high-reliability. The commercial inspections are given to chips and packages destined for use in commercial applications. The high-reliability specifications are derived from a set of government standards identified as “Mil-Standard 883.” Commercial-level inspections screen the bonded chips for die-attach quality, correct placement of the bonds on the bonding pads and inner leads, the shape and quality of the ball or wedge bond, and the general condition of the chip surface in regard to contamination, scratches, *etc.* Mil-Standard 883 covers the same general issues as the commercial inspections but to more stringent requirements. In particular, this standard also specifies criteria for the chip surface, including pattern alignment, critical dimensions, and surface irregularities, such as small scratches, voids, and small defects. These tighter criteria address increased reliability in the rigorous conditions encountered in space and military operations.

Sealing Techniques

After the bonded chip passes the optical inspection, it is ready for sealing in a protective enclosure. Several methods are used to achieve the enclosure of the chip. The method chosen depends on whether the seal must be hermetic or nonhermetic and which type of package is to be used. The principal sealing methods use welded seals, soldered seals, glass-sealed lids, Ceramic Dual In-Line Package (CERDIP) package construction, molded epoxy enclosures, and a sealing layer (see “Bare Die Techniques and Blob Top”) directly on a die bonded directly to a flipchip package or PCB ([Fig. 18.27](#)).



FIGURE 18.27 Package sealing methods.

Metal Can

If the package is a metal-can type, a hermetic seal is achieved by welding the flanged lid to a matching flange on the base of the package.

Premade Packages

Premade ceramic packages are sealed by one of two methods, metal or ceramic lids. Packages made for metal lids have a ring of gold around the top of the die-attach cavity, called the *seal ring*. Placed on top of the seal-ring area is a preformed piece of gold-tin solder. The metal lid is clamped in position over the seal ring and placed on the belt of a conveyor furnace. As the clamped package passes through the furnace, the lid and package are soldered together to form a hermetic seal. The sealing takes place at a temperature range of 320 to 360°C in a pure nitrogen atmosphere. If the package is to receive a ceramic lid, the procedure is similar. The part of the ceramic lid that contacts the base, outside the die-attach area, is coated with a layer of low-melting-point glass. The hermetic seal is completed as the package passes through a conveyor furnace. This sealing takes place at a temperature of about 400°C in an atmosphere of clean, dry air.

CERDIP Packages

A completed CERDIP package results in a hermetic seal around the chip and bonding wires. This seal is accomplished with glass, similar to the ceramic seal on the premade packages. In the case of the CERDIP package, the inner metal-lead system is buried in a glass layer. The ceramic top of the package system has a cavity (Fig. 18.28). The underside of the lid, outside the cavity area, is coated with a layer of low-melting-point glass. The lid is placed over the base and clamped. The seal is accomplished as the assembly is passed through a conveyor furnace or placed in an oven. In the furnace or oven, the glass melts, fusing the base and top together. The CERDIP glass-sealing system is used to seal DIP and flat packs. These latter packages are known as *Cerpacks* and *Cerflats*.

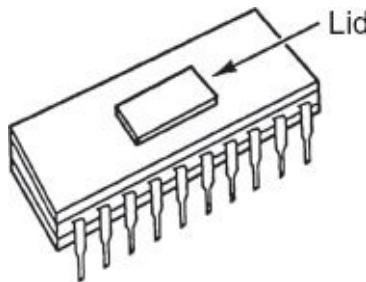


FIGURE 18.28 Premade ceramic package.

Molded Epoxy Enclosures

The fourth major method of enclosure, *epoxy molding*, produces the traditional plastic package ([Fig. 18.29](#)). The resultant seal, while protecting the die from moisture and contamination, is not classified as hermetic. However, there exists a considerable amount of research into the development of improved epoxy materials to create better enclosures. The major advantages of epoxy-molded enclosures are weight, low material cost, and manufacturing efficiency.

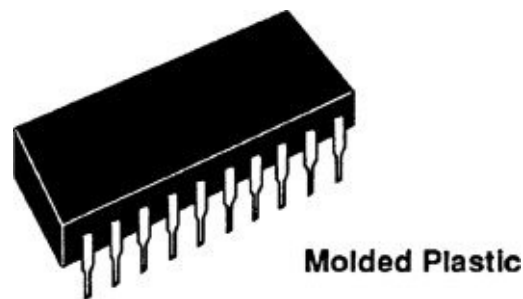


FIGURE 18.29 Molded CERDIP.

This sealing method follows a different process flow. The die is attached and bonded to a lead frame containing a number of lead systems ([Fig. 18.30](#)). After the preseat inspection, the lead frames are transferred to the molding area. The frames are placed on a mold mounted in a transfer molding machine. The molding machine is in turn charged with pellets of the epoxy material, which have been previously softened by a radio frequency heater. Inside, the pellets are forced by a ram into a liquid state. The ram then forces the liquid around the die on the lead frames, forming an individual package around each lead frame. After the epoxy sets in the mold, the frames are removed and placed in an oven for final curing.



FIGURE 18.30 Lead frame.

Lead Plating

An important feature of the completed packages described is the finish on the package leads. Most package leads are coated with lead-tin solder, tin plate, or gold plate. The plating serves several important functions.

First is the solderability of the leads into a circuit board. The additional metal finish improves the lead solderability, resulting in a more reliable electrical connection of the package and the printed circuit.

The second benefit of the lead finish is that it protects the leads from oxidation or corrosion during periods of storage prior to mounting on the circuit board. The third benefit of lead plating is the protection of the leads from corrosive agents in the packaging and printed circuit-board mounting processes. These agents include solder flux, corrosive cleaners, and even tap water. The plating continues to protect the leads during their lifetime of use.

Electrolytic Plating

Plated layers, such as gold and tin, are applied by electrolytic procedures. The packages are mounted on racks with each lead connected to an electric potential. The racks are placed in a tank containing a plating solution. Next, a small current is passed between the packages and an electrode in the tank. The current causes the particular metal in the solution to plate onto the leads.

Tin-Lead Solder

Tin-lead solder layers are applied either by dipping the packages into a pot of the molten solder or by a wave-soldering technique. This latter technique offers the advantage of good control of the layer thickness and provides a shorter exposure of the package to the molten solder.

Plating Process Flows

Metal cans, side-brazed DIP packs, and pin-grid arrays have their leads plated before starting into the packaging process. Cerdip and plastic packages go through the plating process after the sealing steps.

Lead Trimming

One of the last steps in the package-assembly process is trimming away excess material from the leads. The outer leads of DIP and flat-pack packages are made with a tie-bar ([Fig. 18.31](#)). This bar keeps the leads from becoming bent during the packaging process. At the end of the process, the package goes through a simple trimming machine that simultaneously trims away the tie-bar and trims the leads to the proper length.

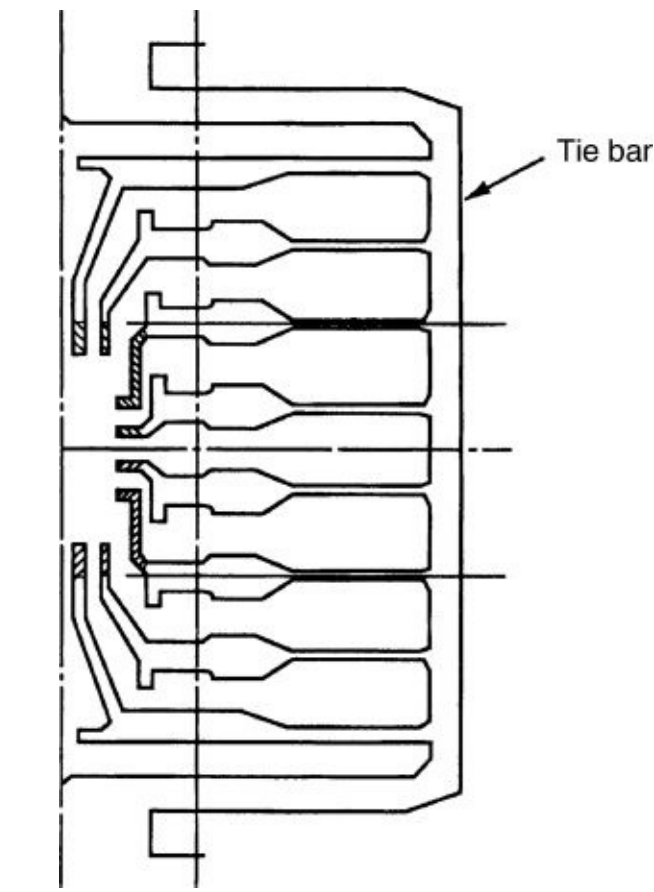


FIGURE 18.31 Tie-bar.

Plastic package lead frames have an additional piece of material. It is a bridge of metal close to the package body that functions as a dam to prevent the liquid epoxy material from running into the lead area ([Fig. 18.32](#)). The dam is cut away from the frame with a series of precise cutting tools. After the dam is cut away, the packages move to another station on the cutter where the frame is separated into individual packages. If the package is a surface-mount type, the leads will be bent into the required shape.

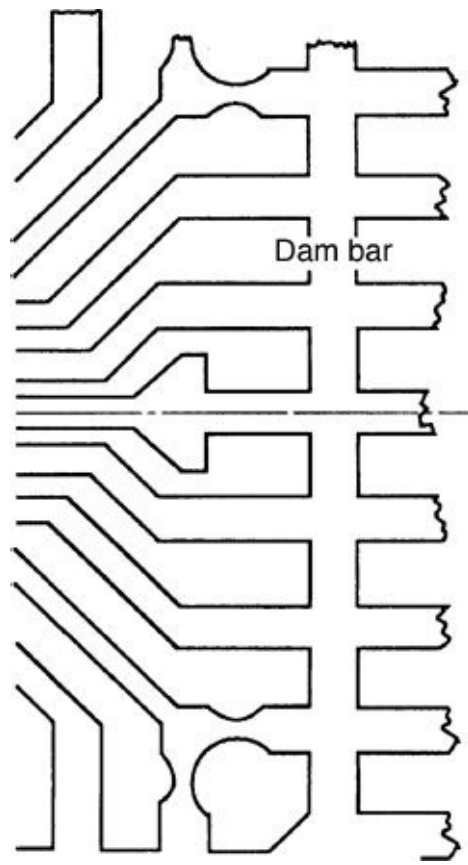


FIGURE 18.32 Lead frame dam.

Deflashing

Plastic packages receive an additional process, called *deflashing*, which is required to remove excess molding material from the package enclosure. Deflashing is done by either dipping the packages in a chemical bath followed by a rinse or by a physical abrasion process. Physical deflashing is done in a machine similar to a sandblaster that uses plastic beads as the abrasive.

Package Marking

Once completed, a package is identified with key information. Typical information coded on the package is the product type, device specifications, date, lot number, and where it was made. The main methods of marking are ink printing and laser inscription. Ink printing has the advantage of good adherence to all of the package materials.

The composition of the ink is chosen for permanence in the eventual operating environment of the device. The ink is applied by an offset printer followed by a curing step. Curing is done by oven drying, room-temperature air drying, or by ultraviolet light. Laser marking is especially well-suited to plastic packages. The mark is permanently inscribed in the package surface and can provide good contrast on the dark packages. Additionally, laser marking is fast and noncontaminating, since no foreign material is added to the package surface, and no curing step is required. A drawback to laser marking is the difficulty of changing the mark if a wrong code was used or the status of the device is changed. Regardless of the marking method, all marks must meet the requirements of legibility (especially on smaller packages) and permanence when exposed to harsh environments.

Final Testing

At the conclusion of the packaging process, the completed package is put through a series of environmental, electrical, and reliability tests. These tests vary in type and specifications, depending on the customer and use of the packaged devices. The tests may be performed on all of the packages in a lot or on selected samples.

Environmental Tests

Environmental tests are performed to weed out leaking and defective packages. Defects detected are loose chips, contaminants and particles in the die-attach cavity, and faulty bonding. This testing series starts with a stabilization bake to drive off any volatile substances in the package. A typical bake is at 150°C for 24 hours.

The first environmental test is *temperature cycling*. The packages are loaded into a chamber and cycled between two temperature extremes. The number of cycles may reach several hundred. The high and low temperatures of this test vary with the device use. Commercial parts receive a narrower temperature range than hi-rel parts. The hi-rel cycle range is -25 to 125°C. During the cycling, any weakness in the seal, die attachment, or bonding will be aggravated and detected in later electrical tests.

A second environmental test is *constant acceleration*. In this test, the packages are accelerated in a centrifuge ([Fig. 18.33](#)) that creates a force as high as 30,000 times the pull of gravity on the Earth (30,000 *g*'s). During the acceleration, loose particles in the package, poorly attached dies, and weakly attached bonds are stressed so that they will be detected at the final electrical tests.

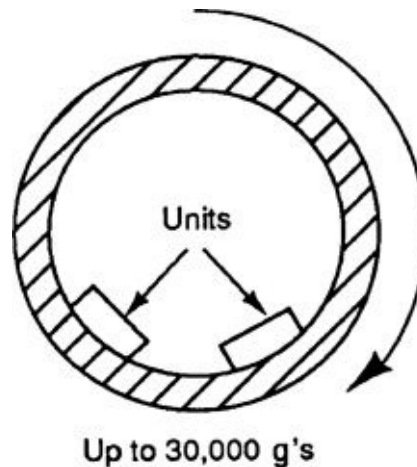


FIGURE 18.33 Acceleration test.

Leaks in the package enclosure are detected by two tests. *Gross leak* testing ([Fig. 18.34](#)) is conducted by submerging the packages in hot liquid. The heated liquid raises the temperature of the package and forces trapped gases in the cavity to escape. The escaping gases are observable as bubbles rising in the liquid. The chamber has a transparent side, allowing an operator to observe the bubbles. Smaller (or fine) leaks are detected by using tracer gases. For this check, helium is pumped under pressure into a chamber containing the packages. If the package has

small leaks, the gas will be pumped into the package cavity. The gas is detected as it escapes through the small leaks by a machine known as a mass spectrometer, which can identify the escaping gas. An alternate fine leak test uses the radioactive gas krypton-85. It too is pumped through any leaks into the package under pressure. Detection of any krypton-85 in the package is by a device similar to a Geiger counter.

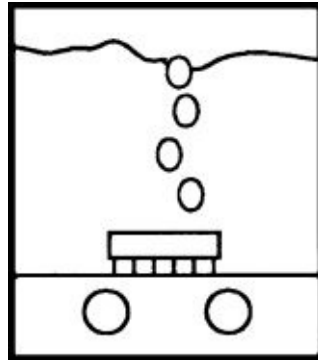


FIGURE 18.34 Gross leak bubbles.

Electrical Testing

The purpose of the wafer-fabrication and -packaging processes is to provide to the customer a specific semiconductor device that performs to specific parameters. Thus, one of the last steps is an electrical test of the completed unit to verify that it performs to specifications. The tests are similar to the wafer-sort tests. The overall objective is to verify that the good chips identified at wafer sort have not been compromised by the packaging process.

First, there is a series of parametric tests. These electrical tests check the general performance of the device or circuit and ensure that it meets certain input and output voltage, capacitance, and current specifications. The second part of the final test is called the *functional test*. This test actually exercises the specific chip functioning. Logic chips are put through logic tests, and memory chips are exercised in their data-storage and -retrieval capabilities. The equipment used to conduct the final test is electrically similar to that used in the wafer-sort operation. The electrical tests are performed by a computer-controlled tester that directs the sequences and levels of the parametric and functional tests. The packages are connected to the tester through sockets; the socket unit is known as the *test head*. The packages are inserted into the test head manually or by an automatic unit known as a *handler* ([Fig. 18.35](#)). This handler may be mechanical or robotic, depending on the speed and complexity of the operation.

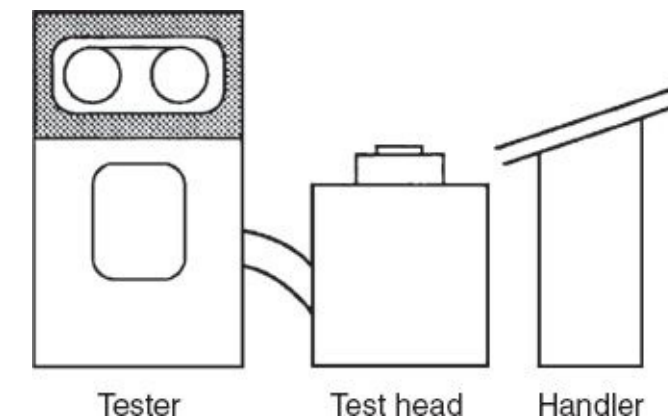


FIGURE 18.35 Final test.

Burn-In Tests

The last of the tests is the optional burn-in test(s). The reason it is optional is that, although it is required for all high-reliability device lots, it may or may not be performed on commercial devices. The test requires the insertion of the packages in sockets and mounting in a chamber with temperature-cycling capability. During the

test, the circuits are temperature-cycled while under an electrical bias.

The burn-in test is intended to stress the electrical interconnection of the chip and package and drive any contaminants in the body of the chip into the active circuitry, thus causing failure. This test is based on data that indicate that chips prone to these types of failures actually malfunction in the early part of their lifetime. By conducting burn-in tests, the early failures are detected. The devices passing the test are statistically more reliable.

Package Design

Up to the early 1970s, most chips ended up in either a metal package, known as a “can,” or in the familiar dual in-line package (DIP). The trends in chip size and integration levels and new electronic products with special packaging requirements (handheld mobile devices) have driven the development of new packages and strategies. Certainly bonding techniques are a major driver of package design. Function and component density are also major drivers. ICs evolved from specific functions (logic, memory) to microprocessors and into entire systems on a chip (SoC). Packages are also undergoing a similar evolution with packages being designed for higher level functioning all in one package.

A family tree of package design is in [Figure 18.36](#). On the single-chip side, there are a seemingly endless list of package types. The more familiar wire-bonded types and the basic flipchip packaging are described in the following sections. On the other side, there is the system-in package (SIP) schemes with the older multichip modules (MCM) and single packages containing an SoC or the 3D package containing several chips in a vertical stack.

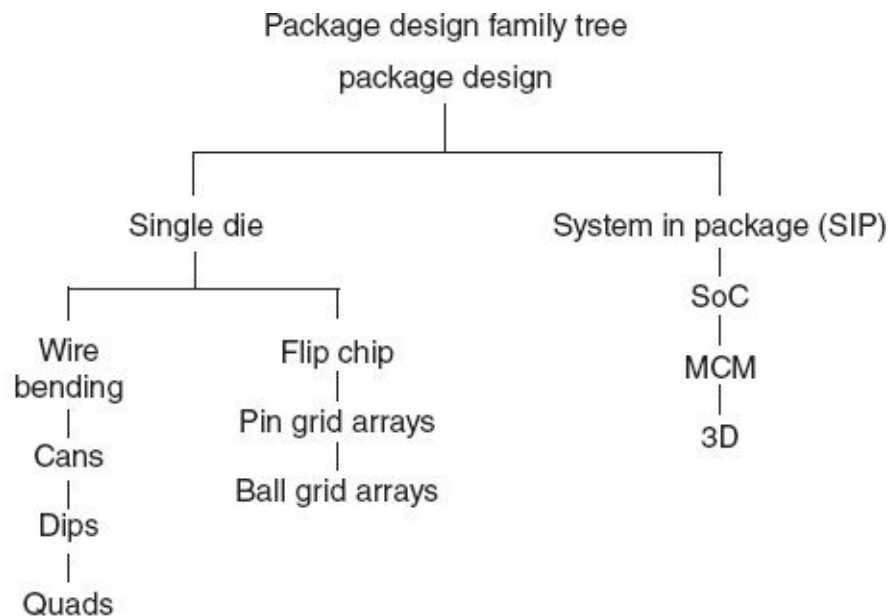


FIGURE 18.36 Package design family tree.

Metal Cans and Dual In-Line Packages

Metal cans are cylindrically shaped packages with an array of leads extending through the base ([Fig. 18.37](#)). The chip is attached to the base and wire bonded to posts that are connected to the leads. The lid and base have matching flanges that are

welded together to create a hermetic seal. These packages are designated by numbers, with the T0-3 and T0-5 being the most common. Metal cans are used to package discrete devices and small-scale integrated circuits.

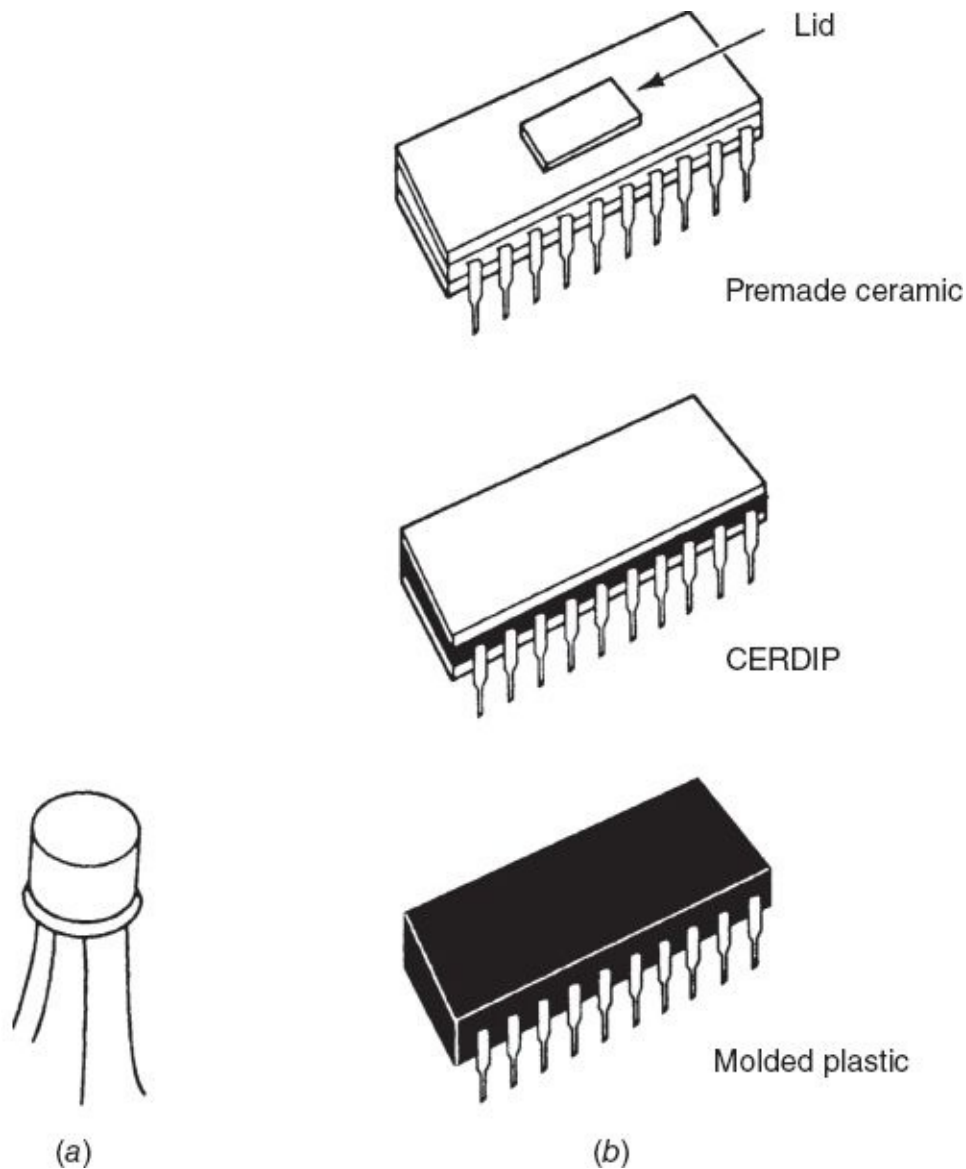


FIGURE 18.37 (a) Metal can and (b) DIP packages.

The DIP is probably the most familiar package design. It features a thick sturdy body with two rows of outer leads coming out of the side and bending downward. DIPs are constructed by three different techniques ([Fig. 18.38](#)). Chips designed for high-reliability use will be packaged in a premade ceramic DIP. The package is formed with a solid body of ceramic with the leads buried in the ceramic. The die-attachment area is a cavity recessed into the body. The hermetic seal is completed by a soldered metal lid or a glass-sealed ceramic lid.

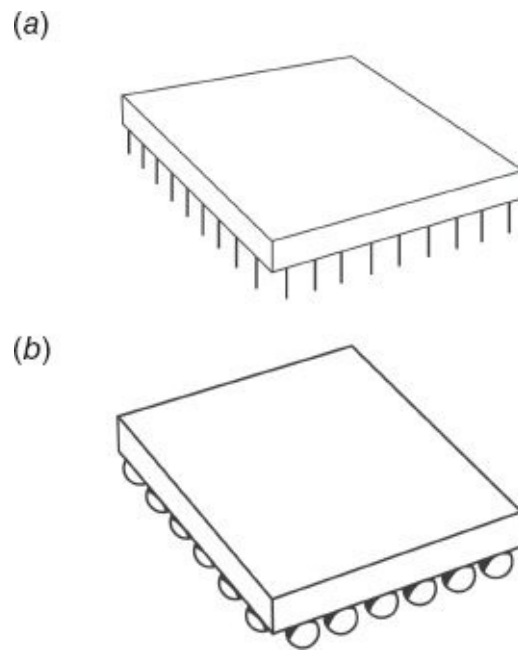


FIGURE 18.38 (a) Pin grid array, and (b) ball grid array.

Another approach to the DIP is the Cerdip, which stands for ceramic DIP. This type of package is composed of a bottom ceramic base with the lead-frame held firmly in a glass layer. The chip is attached to the base and wire bonded to the lead frame. The hermetic seal is completed with a ceramic top glass sealed to the base. Cerdip construction is used for a number of package types. The vast majority of DIPs are made by the epoxy-molding technique. In this technique, the chip is attached to a lead frame and then wire bonded. After bonding, the frame is placed in a molding machine and the package is formed around the chip, wire bonds, and inner leads.

Pin Grid Arrays

Larger chips, with more leads, have outgrown the DIP configuration. The pin grid array is a package designed for larger chips. It features a premade “sandwich” with the outer leads coming out of the bottom of the package in the form of pins ([Fig. 18.38](#)). The chip is attached in a cavity that is formed in either the top or bottom of the body, usually using bump or flipchip technology. Connections to the chip cover the entire chip area, unlike most chips with connections restricted to bonding pads around the periphery of the chip. Ceramic pin grid arrays (PGAs) are hermetically sealed with a soldered metal lid.

Ball-Grid Arrays or FlipChip Ball-Grid Arrays

Ball-grid arrays have the same body shape as PGAs. Instead of pins on the package bottom, there is a series of solder bumps (balls) used to connect the package to the PCB ([Fig. 18.38b](#)). This is essentially the same bump or ball technology used to connect die to packages.

The effect is to lower the package profile and weight as well as providing higher pin counts by using the whole-chip surface for bonding pads. Balls (or bumps) also bring the aspect of absorbing stresses created from the thermal expansion differences between the package and the PCB.

Quad Packages

While pin grid array packs are a convenient design for larger chips, their ceramic construction is expensive compared to molded epoxy packages. This consideration led to the development of the “quad” package. A quad (short for quadrant) pack is constructed by the epoxy-molded technique but has leads coming out of all four sides of the package ([Fig. 18.39](#)) allowing a surface-mount technology (SMT) for lower profiles on PCBs. The chip-to-package bonding can be wire bonding or BGA schemes.

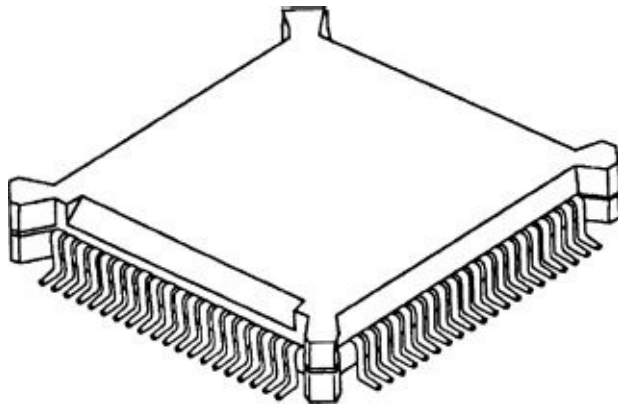


FIGURE 18.39 Quad package.

Thin Packages

Several techniques are used to make thinner packages. They are called flat packs (FPs), thin small outline packages (TSOPs), small outline IC (SOIC), or ultra-thin packages (UTPs). Flat packs are constructed by the same techniques used to form DIP packages. These packages are designed with flatter height profiles and have their leads bent out to the side of the package ([Fig. 18.40](#)). Ultrathin packages have total body heights in the 1-mm range. There are also quad flat packs (QFPs).

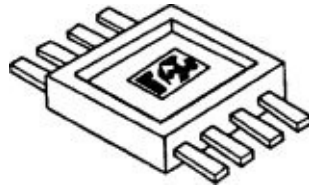


FIGURE 18.40 Flat pack.

Chip-Scale Packages

In the world of ICs, the perfect package is no package. It is recognized that any package brings with it electrical resistance, weight, the opportunity to degrade the circuit performance, and cost. Overall, the smaller the package, the cheaper the cost of packaging and the benefit of higher densities. Chip-scale packages meet this need ([Fig. 18.41](#)). They are simply packages with dimensions within 1.2 times the die size.⁸ The challenges are to provide adequate mechanical and environmental protection for easy connection to PCBs.

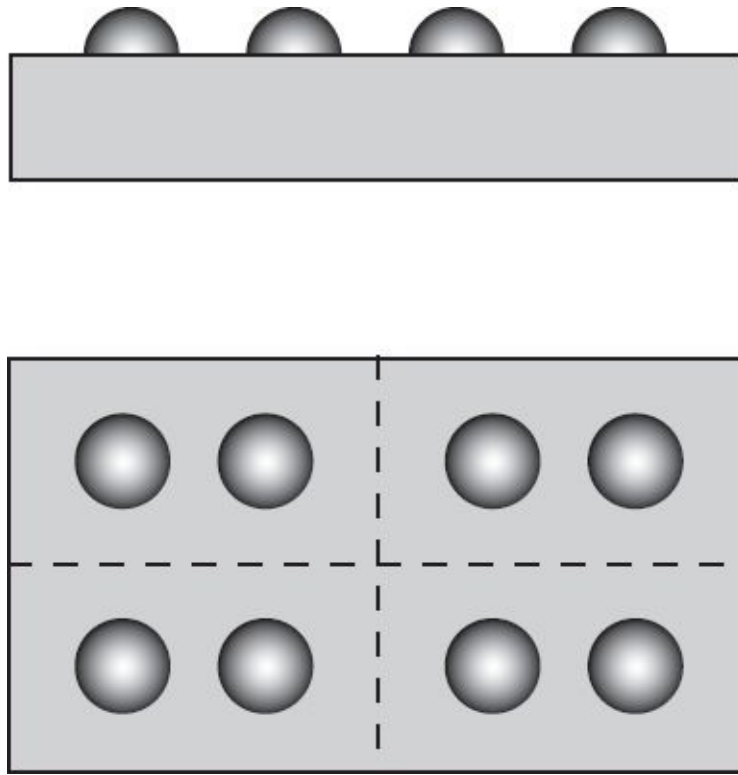


FIGURE 18.41 Chip-scale package.

General design approaches for chip-scale packaging favor flipchip technology with ball grid arrays and blob top protection. The march to smaller packages and more reliable electrical connections has led to the micro-ball grid array, or μ BGA.

Lead on Chip

The LOC package, intended for large die, has the bonding pads arranged down the middle of the chip. Package leads sit on a cushion over the chip surface.

Three-Dimensional Packages

“Beyond Moore” has become a catch phrase in packaging literature. It notes that IC densities are maxing out as transistor scaling is reaching physical limits. It also recognizes that “beyond Moore” means technologies that pack more functions in a package under the general term of system in package (SIP). The industry is developing numerous approaches with two basic approaches: stacking die and stacking packages [package on package (POP)].

Stacking Die Techniques

There are four die-stacking technologies: Monolithic, wafer on wafer, die on wafer, and die on die.

Monolithic

The monolithic technology is to build multiple circuit layers in the wafer-fabrication process. The individual circuit layers are interconnected by multiple metal layers using plug or via techniques.

Wafer-on-Wafer

Individual wafers with different circuits are thinned, bump or ball bonded to each other. 3-D separation leaves a three-dimensional stack all interconnected. Some systems utilize vias drilled or etched through the wafers to allow bonding. These are called through-silicon vias (TSV). Other arrangements have different sized die wire bonded to the package substrate. Another scheme is to bump or ball bond individual wafers and also use wire bonding for additional connections in the package ([Fig. 18.42](#)).

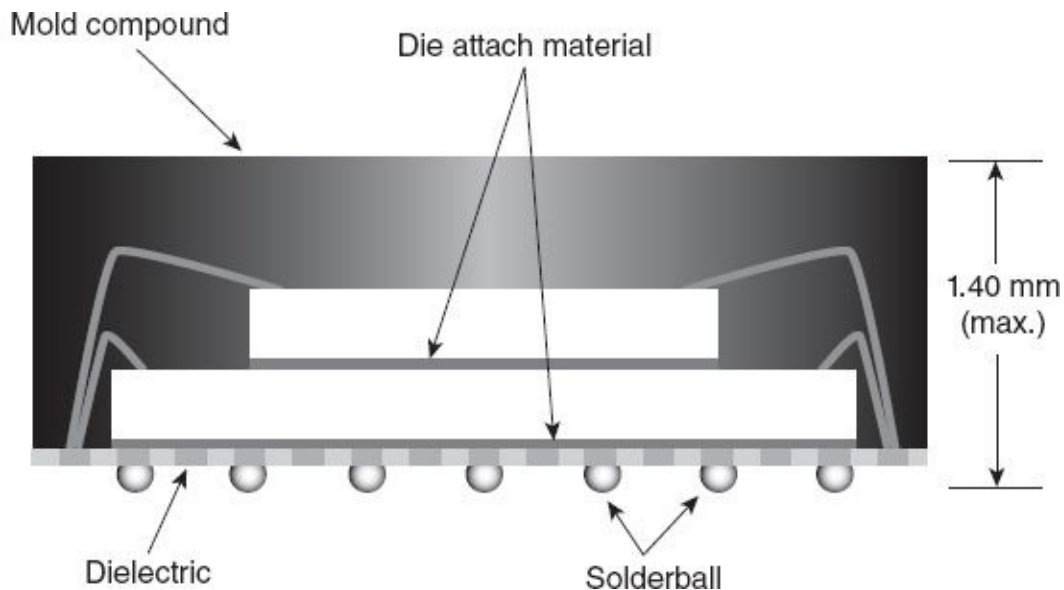


FIGURE 18.42 Example stacked die.

Die-on-Wafer

Die from one wafer are bonded onto sites on chips of another wafer before die separation. Connections are through TSVs and bump or ball bonding.

Die-on-Die

Die from separate wafers are diced and connected through TSVs and bump/ball bonding into an integrated stack. An advantage of this scheme is higher package yields because the individual die are known-good-die going into this packaging

process.

Package on Package or Package in Package

As the name implies, packaging individual die (or stacked die) and then stacking the packages also achieves higher functioning “systems” for a given area. Again, combinations of wire bonding and bump or ball bonding may be necessary¹⁰ ([Fig. 18.43](#)).

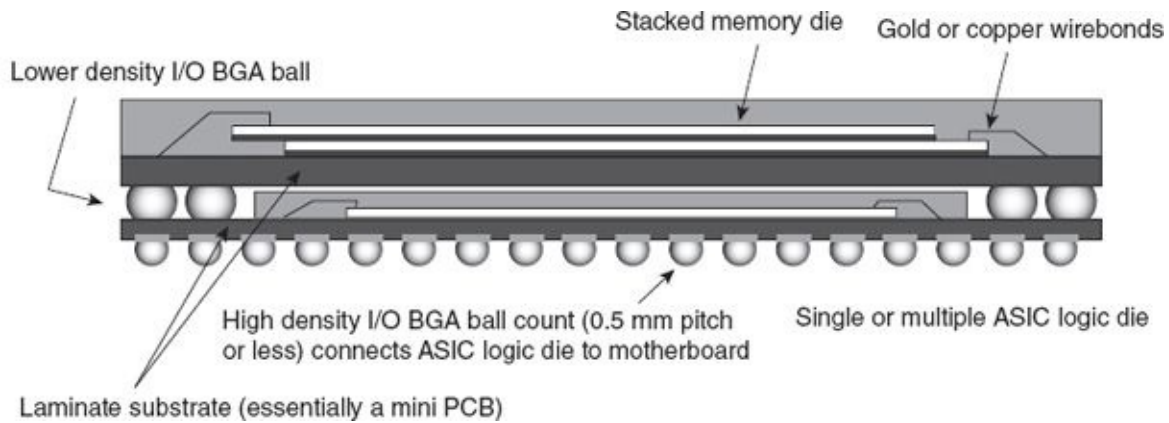


FIGURE 18.43 Example package on package design.

Through Silicon Vias

Besides face-to-face bump or ball bonding of die and wire bonding from top die to lower die (dice) there is *through silicon vias* (TSVs). The concept is simple, create top to bottom holes (vias) in the wafer. Both mask or etch and plasma-etching are used. The vias can be created in the front end (FEOL) of the process or at the back end (BEOL) of the wafer-fabrication process ([Fig. 18.44](#)).

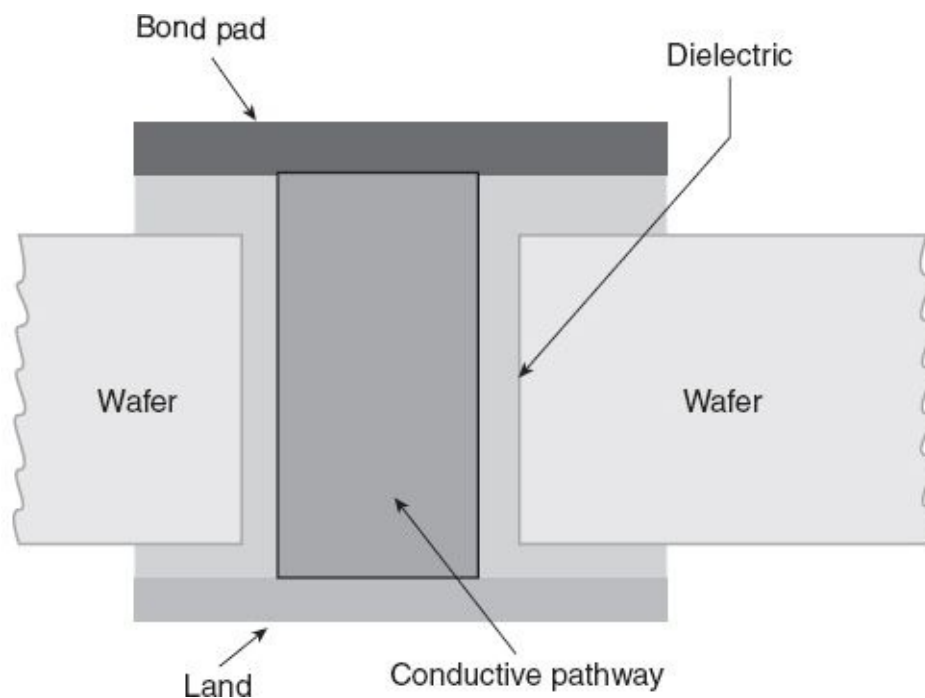


FIGURE 18.44 Through silicon via schematic. (Source: Tessara.)

Interposers

Interposers are intermediate laminate layers or substrates that allow an additional connection plane (or planes) between die in the package and/or between die and package connections (Fig. 18.45). They also can supply rigidity to the package. Organic interposers are made of layers of laminates with through-hole conductors. Rigid interposers are generally made from epoxy materials while flexible interposers are made from resins.¹¹ In some schemes, a silicon slice with through silicon vias (TSVs) are used as interposers.

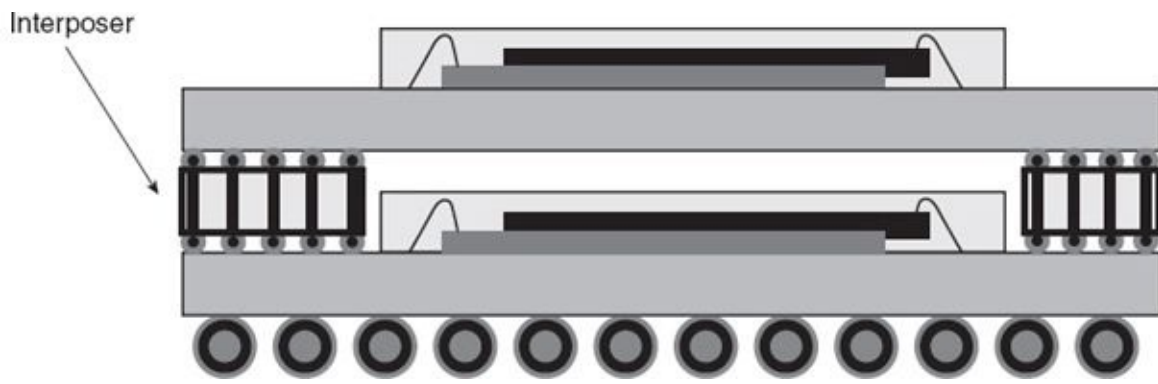


FIGURE 18.45 Example of an interposer in POP system.

Three-Dimensional Enabling Technologies

The goal of stacking multiple die in a package while keeping the height thin and the foot print small brings together many advanced packaging technologies. In package die-to-die bonding is by bump or ball technology is applied to the wafer at the end of the traditional wafer-fabrication process (See “Wafer Scale Packaging”). Next is the thinning the wafers with CMP and a polishing process down to the sub 100- μm range. At this thickness wafers are prone to breaking and curling. The solution is to attach a thin die *attach film* to the wafer back prior to die separation by sawing. Generally, the bottom die is attached to the package by epoxies to accommodate any unevenness in the package die attach area.

Hybrid Circuits

Hybrid circuitry is an old technology, long favored for use in military and harsh environments. The term *hybrid* refers to the mix of semiconductor and conventional passive (resistors, capacitors) electrical devices present in the same circuit. The devices are connected by conductive or resistive thick-film paths on the surface of an insulating substrate. These paths are formed by silk screening inks containing a proper filler or from thin films evaporated or sputtered onto the surface and patterned by photolithography techniques. Most hybrid substrates are ceramic. High-performance hybrids may have AlN, SiC, or diamond substrates used alone or in combinations.^{[12](#)}

Hybrid circuits offer the advantages of structural rigidity and low leakage between devices due to the insulating property of the ceramic. They can provide a circuit with a mix of CMOS, bipolar and other components and offer functions not available in ASIC circuits.^{[13](#)}

On the downside, they generally have a much lower density than packaged integrated circuits and have a higher cost.

Multichip Modules

Mounting individual chip packages on a PCB present several challenges. A chip package is several times the area of the chip taking up space on the board. Circuit resistance is increased by the individual resistances of all the package pins, and the longer cumulative electrical path from the chip bonding pads to the connections to the circuit board. Each of these problems is reduced by mounting several chips on the same substrate. MCMs make use of interposer along with bumps and wire bonding to vertically and horizontally connect the various chips¹⁴ ([Fig. 18.46](#)).

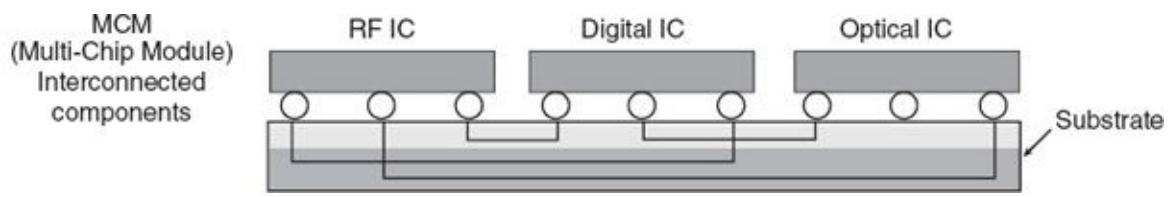


FIGURE 18.46 MCM cross-section.

The Known Good Die Problem

In the individual package process, a final test assures quality of the completed product. If the chip has gone bad or the process was faulty, the entire chip and package is discarded. But the cost is much higher when putting bare die into hybrid, MCM, and PCBs. These vehicles are more expensive to build and carry other expensive chips or components.

One option is to rely on the results of the wafer-sort test to certify die performance. Unfortunately, wafer sort does not include environmental tests or long-term reliability tests. However, performing a final test on a bare die is difficult if not impossible.

Package Type or Technology Summary

There are thousands of individual package types, and there is no uniform system of identifying them. Some are named by their bonding technology, design (DIP, flat packs, and so on), some are named by their construction technique (molded, CERDIP, and so on), and others are named by their use. When trying to understand a package type, keep in mind the three considerations: bonding, design type, construction technique, and use.

In the near future flipchip ball-grid arrays, and chip-scale packages (CSPs), multichip modules (MCMs), and 3D packaging schemes will continue to play a big role in achieving increased system functions.

Package or PCB Connections

On the package side, there are four primary techniques used to connect the package to a PCB.¹⁵ They are through-hole, surface mount, TAB, and ball (bump) technology. Through-hole connections feature straight pins on the package, which are inserted into corresponding holes in the PCB ([Fig. 18.47](#)). *Surface mount* packages have their leads bent into a J shape or bent outward to allow direct soldering to the board surface ([Fig. 18.48](#)). Some surface mount device (SMD) packages do not have leads, rather, they terminate in metal traces hugging the package body. They are known as *leadless packs*. For inclusion on a circuit board, they are inserted into chip carriers, which in turn have the leads that attach to the PCB. Last, there is the bump technology used to connect die to packages, adapted to connecting packages to a PCB. Tape automated bonding (TAB) has two uses. One is bonding the chip bonding pads directly to the lead frame (see section on bonding). TAB is also a technique for bonding the outer leads directly to the PCB.

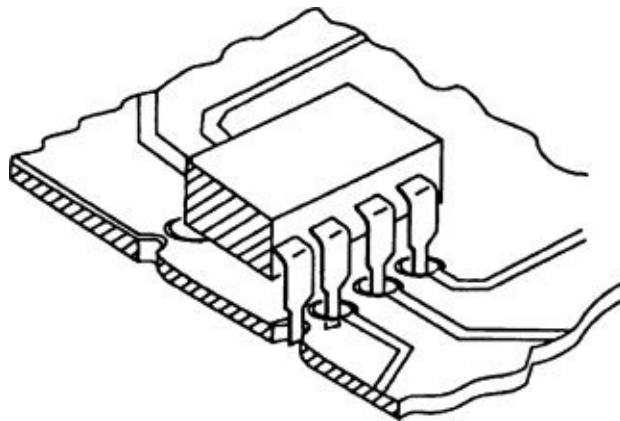


FIGURE 18.47 DIP through-hole assembly.

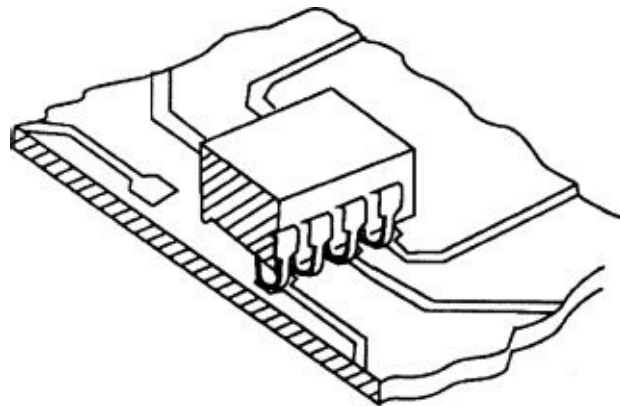


FIGURE 18.48 Surface-mount device.

Bare Die Techniques and Blob Top

Increased reliability with higher density and faster circuits presents an ongoing goal and challenge. Reliability is addressed in hybrid circuits. Speed and higher density comes with the elimination of the individual die package. Fewer links in the chain reduce resistances and (in some cases) shorten the length carriers have to travel, increasing speed. This strategy, called *bare die strategy*, is used in hybrid, multichip modules, and chip on board technology.

The most direct use of bare die is to bond them directly to the PCB. Bonding techniques include all the ones used to bond chips into packages. Protecting the chip after connection to the board is by *blob-top protection*. The protection is accomplished by covering the chip and bonds with a blob of epoxy resin material ([Fig. 18.49](#)).

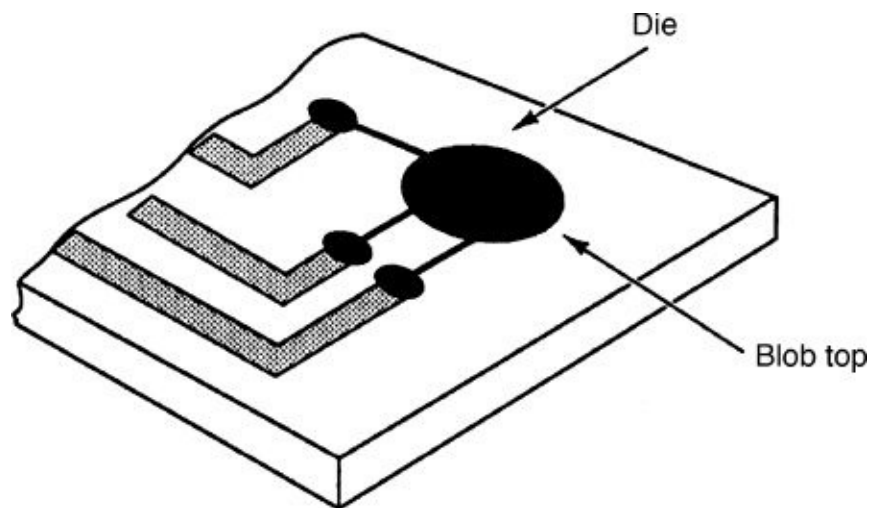


FIGURE 18.49 Blob top.

The material is similar in properties to that used to mold the plastic packages. Blob-top coverings are used with TAB and other packaging schemes.

Review Topics

Upon completion of this chapter, you should be able to:

1. List the four functions of a semiconductor package.
 2. List the five common parts of a package.
 3. Recognize and identify the major package designs.
 4. List and describe the major packaging process flows.
-

References

1. Data sheet, CORWIL Technology Corporation, 2007.
2. Blech, F. I., and Dang, D., “Silicon Wafer Deformation after Backside Grinding,” *Solid State Technology*, PennWell Publishing, Aug. 1994:74.
3. Plummer, L., “Packaging Trends,” *Semiconductor International*, Cahners Publishing, Jan. 1993:33.
4. Iscoff, R., “Ultrathin Packages: Are They Ahead of Their Time?” *Semiconductor International*, Cahners Publishing, May 1994:50.
5. Tummala, R., and Rymaszewski, E., *Microelectronics Packaging Handbook*, 1989, Van Nostrand Reinhold, New York, NY.
6. Karnezos, M., *3-D Packaging: Where all the Technologies Come Together*, IEEE/Semi Int’l Electronics Manufacturing Technology Symposium, 2004.
7. Riley, G., *Under Bump Metalization (UBM)*, <http://flipchips.com/tutorial/process/under-bump-metallization-ubm>, (Accessed on: Sep. 2001).
8. DiStefano, T., and Fjelstad, J., “Chip-Scale Packaging Meets Future Design Needs,” *Solid State Technology*, Apr. 1996:82.
9. Kada, M., *Advancements in Stacked Chip Scale Packaging (S-CSP)*, Proceedings of Pan Pacific Microelectronics Symposium Conference, Jan. 2000.
10. Cunningham, A., *The PC Onside Your Phone—A Guide to the System-on-a-Chip*, Arstechnica.com, (Accessed on Apr. 10, 2013).
11. Rabindra, N., Das, R. N., Egitto, D., Bonitz, B., et al., *Markovich Package-Interposer-Package (PIP) Technology for High End Electronics*, Endicott Interconnect Technologies, Inc., New York, NY:13760.
12. Rao, Tummala, R. R., “System on System Integrates Multiple Tasks,” *The International Magazine for Device and Packaging*, Feb. 2004:101.
13. Nguyen, N., “Using Advanced Substrate Materials with Hybrid Packaging Techniques for Ultrahigh-Power ICs,” *Solid State Technology*, PennWell Publishing, Feb. 1993:59.
14. Iscoff, R., “Will Hybrid Circuits Survive?” *Semiconductor International*, Cahners Publishing, Oct. 1993:57.
15. Baliga, J., “Package Styles Drive Advancements in Die Bonding,” *Semiconductor International*, Jun. 1997:101.

Glossary

3D packaging Stacking and connecting two or more die in a single package.

III-V semiconductor materials Semiconducting materials formed by elements from columns III and V of the periodic table of the elements.

II-VI semiconductor materials Semiconducting materials formed by elements from columns II and VI of the periodic table of the elements.

Acceptor An impurity that causes semiconducting materials to accept valence electrons, thereby leaving “holes” in the valence band. The holes act like carriers of positive charge, referred to as P-type.

Airborne molecular contamination (AMC) Contaminating airborne molecules that are present in clean room air.

Aligner (align and expose) A process tool used to align wafers and masks or reticles and expose the photoresist with a UV or other radiation source.

Alignment Refers to the positioning of a mask or reticle with respect to the wafer.

Alignment marks Targets on the mask and wafer used for correct alignment.

Alloy (1) A compound composed of two metals. (2) In semiconductor processing, the alloy step causes the interdiffusion of the semiconductor and the material on top of it, forming an ohmic contact between them.

Aluminum (Al) The metal most often used in semiconductor technology to form the interconnects between devices on a chip. It can be applied by evaporation or sputtering.

Amorphous Materials with no definite arrangement of atoms, e.g., plastics are amorphous.

Amplified resist A photoresist whose chemical reactions are enhanced with added chemicals.

Angstrom A unit of length, an angstrom (Å) is one ten-thousandth of a micron (10^{-4} μm). 100,000,000 (i.e., 100 million) Å = 1 cm.

Anisotropic An etch process that exhibits little or no undercutting.

Anneal A high-temperature processing step (usually the last one), designed to minimize stress in the crystal structure of the wafer.

Antimony (Sb) A Group V element that is an N-type dopant in silicon. It is often used as the dopant for the buried layer.

Antireflective coating (ARC) A chemical layer added to a wafer surface to reduce reflections during exposure.

Arsenic A Group V element that is an N-type dopant in silicon.

Assembly The series of operations after fabrication in which the wafer is separated into individual chips and mounted and connected to a package.

Atmospheric oxidation A process of oxidation of silicon carried out at atmospheric pressure. The equipment used for thermal oxidation is the same as that used for thermal diffusion.

Atomic force microscope (AFM) A microscope for profiling wafer surfaces by plotting the output of a spring-balanced probe moved over the surface.

Atomic layer deposition (ALD) A method to build up (deposit) a layer one atom layer at a time.

Atomic number A number assigned to each element, equal to the number of protons (therefore also the number of electrons) in the atom.

Atomic particles The parts of an atom: electrons, protons, and neutrons.

Base (1) The control portion of an NPN or PNP junction transistor. (2) The P-type diffusion done using boron that forms the base of NPN transistors, and the emitter and collector of lateral PNP transistors and resistors.

Bi-MOS A circuit containing both bipolar and MOS transistors.

Binary notation A way of representing any number using a power of 2 (i.e., using the digits 0 and 1).

Bipolar transistor A transistor consisting of an emitter, base, and collector, whose action depends on the injection of minority carriers from the base by the collector.

Sometimes called NPN or PNP transistor to emphasize its layer structure.

Boat (1) Pieces of quartz or metal joined together to form a supporting structure for wafers during high-temperature processing steps. (2) A Teflon® or plastic assemblage used to hold wafers during wet processing steps.

Boat puller A mechanical arrangement to push a boat loaded with wafers into a furnace and/or withdraw it at a fixed speed.

BOE See [*buffered oxide etch*](#).

Bonding pads Electrical terminals on the chip (generally around the periphery) used for connection to the package electrical system.

Boron (B) The P-type dopant commonly used for the isolation and base diffusion in standard bipolar integrated circuit processing.

Boron trichloride (BCl₃) A gas that is often used as a source of boron for doping silicon.

Bubbler An apparatus in which a carrier gas is “bubbled” through a heated liquid, causing portions of the liquid to be transported with the gas, e.g., a carrier gas (nitrogen or oxygen) is bubbled through deionized water at 98 to 99°C on its way to the oxidation tube.

Buffered oxide etch (BOE) A mix of hydrogen fluoride (HF) and ammonium fluoride (NH₄F), used to allow oxide etching to occur at a slow, controlled rate.

Bump/ball connection technology A structure of metal bumps or balls is formed on the bonding pads, allowing connection of the chip to a package when the chip is flipped over.

Buried layer The N⁺ diffusion in the P-type substrate done just prior to growing the epitaxial layer. The buried layer provides a low-resistance path for current flowing in a device. Common buried layer dopants are antimony and arsenic.

Can A metal package used for connecting a chip to a printed circuit board with from three to five leads.

Capacitor A discrete device that stores electrical charge on two conductors that are separated by a dielectric.

Capacitance Electrical charge storage capability.

Capacitance-voltage plot (C/V plot) A plot that provides information on the amount of mobile ionic contamination present in the oxide.

Carrier gas An inert gas that will transport atoms or molecules of a desired substance to a reaction chamber.

Carrier illumination junction detection A nondestructive system of determining a junction depth, from a carrier charge buildup at the junction that is induced by a laser beam.

Centistokes Units used to measure viscosity; centipoise divided by density.

Channel A thin region of a semiconductor that supports conduction. A channel may occur on a surface or in the bulk, essential for the operation of MOSFETs and SIGFETs. In cases where channels are not part of the circuit design, their presence may indicate contamination problems or incomplete isolation.

Channeling A phenomenon in which an ion beam will penetrate into the crystal planes of the wafer. Preventing channeling is accomplished by cutting the wafer “off orientation.” The effect is to tilt the crystal planes relative to the beam direction.

Charge carrier A carrier of electrical charge within the crystal of a solid-state device, such as an electron or hole.

Chemical etching Selective removal of material by means of liquid reactants. The precision of the etch is controlled by the temperature of the etchant, the time of immersion, and the composition of the acid etchant.

Chemical mechanical polishing (CMP) A wafer-flattening and polishing process that combines chemical removal with mechanical buffing. Used for polishing/flattening wafers after crystal growing and wafer planarization during the wafer fabrication process.

Chemical vapor deposition See [CVD](#).

Chip Die or device, one of the individual integrated circuits or discrete devices on a wafer.

Chip scale package A chip package scaled to the chip size.

Chrome A metal often used in mask fabrication to form the layer in which the circuit pattern is generated.

Circuit board See *printed circuit board*.

Circuit layout The calculation of the physical device dimensions required to produce the required electrical parameters. Vertical dimensions determine CVD and doping thickness specifications. Horizontal dimensions determine the wafer pattern dimensions and are the basis for a scaled drawing of the finished circuit (composite drawing).

Class number Number of contaminant particles in a cubic foot of air.

Cleanroom An area in which semiconductor device fabrication takes place. The cleanliness of the room is highly controlled so as to limit the number of contaminants to which the semiconductor is exposed.

Clear field mask A mask on which the pattern is defined by the opaque areas.

Cluster tool Several process stations or tools served by one loading–unloading chamber and wafer-transport system.

CMOS (complementary field-effect transistor) N-and P-channel MOS transistors on the same chip.

Collector Along with the emitter and base, one of the three regions of the bipolar type of transistor.

Collimated light Light in which the rays are parallel; used for gross visual inspection of surfaces.

Composite drawing A scaled drawing of the finished circuit.

Conductivity The ability of materials to conduct electricity (measured in siemens for conductance and in ohms for resistance).

Conductor A material that has low resistivity and high conductivity.

Contact The regions of exposed silicon that are covered during the metallization process to provide electrical access to the devices.

Contact aligner An aligner tool that clamps the wafer and mask into a tight contact

before the resist exposure cycle.

Contact mask The step at which holes are put into the wafer layers to allow the metal layer to reach down to the doped silicon substrate.

Contamination A general term used to describe unwanted material that adversely affects the physical or electrical characteristics of a semiconductor wafer.

Copper (Cu) Metal used to connect semiconductor devices on a chip surface. Generally used with a dual-damascene patterning process.

Critical dimensions (CDs) The widths of the lines and spaces of critical circuit patterns as well as the area of contacts.

Cryogenic pump A vacuum pump that can produce a vacuum to the 10^{-10} torr range, the same level as the vacuum of space. It does not require forepumps or cold traps and is faster than other types of vacuum pumps.

Cryogenic wafer cleaning A wafer surface cleaning technique using a “snow” of high pressure carbon dioxide (CO_2).

Crystal A material in which the atoms are arranged in structured groups called *unit cells*.

Crystal defects Vacancies and dislocations in a crystal that influence the electrical performance of a circuit.

Crystal orientation The orientation of the primary crystal plane, expressed as Miller indices.

Crystal planes The planes in the semiconductor crystal structure along which the die must be aligned so as to prevent “ragged” die edges when the wafer is separated into individual die.

CUM yield See [*fabrication yield*](#).

Current A measure of the number of charged particles passing a given point per unit time.

Curve tracer A piece of electrical test equipment that displays the characteristics of a device visually on a screen.

CVD (chemical vapor deposition) A method for depositing some of the layers that function as dielectrics, conductors, or semiconductors. A chemical containing atoms of the material to be deposited reacts with another chemical, liberating the desired material, which then deposits on the wafer while by-products of the reaction are removed from the reaction chamber.

Czochralski crystal grower A type of crystal grower that uses a seed to pull a crystal from a crucible of molten material.

Dark field mask A mask on which the pattern is defined by the clear portion of the mask.

Deep ultraviolet (DUV) A light wavelength often used to expose photoresist; it has the advantage of an ability to produce smaller image widths.

Defect density The density of defects per square centimeter on a chip.

Dehydration baking A heating process by which wafer surfaces are restored to a hydrophobic condition by baking. Surface water is evaporated from the wafer at elevated temperatures.

Deionized (DI) water Process water that is free of dissolved ions. Specification levels are generally 15 to 18 m Ω of resistance.

Depletion layer The region in a semiconductor where essentially all charge carriers have been swept out by the electric field that exists there.

Deposition Process in which layers are formed as the result of a chemical reaction in which the desired layer material is formed and coats the wafer surface.

Design rule The minimum feature size of a circuit.

Develop inspection The first inspection in the photomasking process, consisting of measurement of critical dimensions and inspection for defects. It is done after development, or after development and hard bake if an automatic baking system is used.

Development A photoresist processing step in which photoresist is removed from areas defined by the masking and exposure step of wafer fabrication.

Developer Chemical used to remove areas defined in the masking and exposure step of wafer fabrication.

Device A single-function component such as a transistor, resistor, or capacitor.

DI water Deionized water; purity of this water is measured by its resistivity, with the standard being 18 MΩ.

Diborane (B₂H₆) A gas that is often used as a source of boron for doping silicon.

Die One unit on a wafer separated by scribe lines; after all of the wafer fabrication steps are completed, die are separated by sawing. The separated units are referred to as *chips*.

Die bonding Assembly step in which individual chips are attached to the package with conductive adhesives or metal alloys.

Die sort See [wafer sort](#).

Dielectric A material that conducts no current when it has a voltage across it. Two dielectrics encountered in semiconductor processing are silicon dioxide and silicon nitride.

Diffusion A process used in semiconductor production that introduces minute amounts of impurities (dopants) into a substrate material such as silicon or germanium and permits the impurity to spread into the substrate. The process is very dependent on temperature and time.

Diffusivity The rate of movement or diffusion of dopants in a semiconductor.

Diode Device that enables current flow in one direction but not in the other.

DIP (dual in-line package) A rectangular circuit package, with leads coming out of the long sides and bent down to fit into a socket.

Discrete device A circuit having a single electrical function. Discrete devices include capacitors, resistors, transistors, diodes, and fuses.

Dislocation A discontinuity in the crystal lattice; a type of crystal defect.

DMOS (diffused MOS) A transistor structure that features a narrow (channel length) separation between the source and drain. The channel length is created by two sequential diffusions through the same hole.

Donor An impurity that can make a semiconductor N-type by donating extra “free”

electrons; electrons carry a negative charge.

Dopant An element that alters the conductivity of a semiconductor by contributing either a hole or an electron to the conduction process. For silicon, the dopants are found in Groups III and V of the periodic table.

Dopant deposition The first step in the doping process, in which the dopant atoms are put into the wafer surface by ion implantation or diffusion.

Doping The introduction of an impurity (dopant) into the crystal lattice of a semiconductor to modify its electronic properties—for example, adding boron to silicon to make the material P type.

Drain Along with the source and gate, one of the three regions of a unipolar or field-effect transistor (FET).

DRAM (dynamic random access memory) Memory device for the storage of digital information. The information is stored in a “volatile” state.

Drive-in Stage in diffusion where the dopant is driven deeper into the wafer.

Dry etch See [*plasma etch*](#).

Dry ox The growth of silicon dioxide using oxygen and hydrogen, which form water vapor at process temperatures, rather than using water vapor directly.

Dry oxide Thermal silicon dioxide grown using oxygen.

Dual damascene A patterning process that first defines the required pattern in a trench in the top wafer surface, followed by overfilling with a conductive metal. The overfill is removed, usually by a chemical-mechanical-polishing process, leaving the metal pattern within the trench.

E-beam (electron beam) An exposure source that allows direct image formation without a mask. An e-beam can be deflected by electrostatic plates and therefore directed to precise locations, resulting in the generation of submicron-size patterns.

E-beam aligner An aligner tool that exposes the resist-coated wafer surface by steering (writing) an electron beam across the wafer surface.

E-beam evaporation (electron beam evaporation) Phase change that uses the energy of a focused electron beam to provide the required energy to change solid

metal or alloys from solid to gas.

E-beam exposure system A machine in which the image pattern is stored in a computer memory and used to control the electrostatic plates that in turn direct the e-beam, resulting in the generation of patterns without the use of reticles or photomasks.

Edge bead A bead that builds up at the edge of the wafer during the photoresist spin process.

Edge die The incomplete die located on the edge of the wafer.

EEPROM (electrically erasable PROM) A memory circuit with the capability of data erasure and acceptance of new information by the application of an electrical pulse.

Electromigration The diffusion of electrons in electric fields set up in the lead while the circuit is in operation. It occurs in aluminum and is exhibited as a field failure, not as a process defect. The metal thins and eventually separates completely, causing an opening in the circuit.

Electron A charged particle revolving around the nucleus of an atom. It can form bonds with electrons from other atoms or be lost, making the atom an ion.

Ellipsometer An instrument that uses laser light sources to measure thin film thickness.

Emitter (1) The region of a transistor that serves as the source or input end for carriers. (2) The N-type diffusion usually done using phosphorus, which forms the emitter of NPN transistors, the base contact of PNP transistors, the N⁺ contact of NPN transistors, and low-value resistors.

Epitaxial (Greek for “arranged upon.”) The growth of a single-crystal semiconductor film upon a single-crystal substrate. The epitaxial layer has the same crystallographic characteristics as the substrate material.

Epoxy package See [*molded package*](#).

EPROM (erasable PROM) A memory circuit with the capability of data erasure and acceptance of new information.

Etch A process for removing material in a specified area through a wet or dry

chemical reaction or by physical removal, such as by sputter etch.

Evaporation A process step that uses heat to change a material (usually a metal or metal alloy) from its solid state to a gaseous state, with the result of the source being deposited on wafers. Both electron beam and filament evaporation are common in semiconductor processing.

Exposure Method of defining patterns by the interaction of light or another form of energy with photoresist that is sensitive to the energy source.

Fabrication Integrated circuit manufacturing processes.

Fabrication yield The percentage of wafers arriving at wafer sort compared with the number started into the process.

Feature size The minimum width of pattern openings or spaces in a device.

FET (field-effect transistor) A transistor consisting of a source, gate, and drain, whose action depends on the flow of majority carriers past the gate from the source to drain. The flow is controlled by the transverse electric field under the gate. See [unipolar transistor](#).

Field oxide The region on an electrical device where the oxide serves the function of a dielectric.

Final test The final assembly step in which the packaged die is put through its last electrical test.

FinFET A 3D transistor design with a built up 'fin' that provides a larger gate area than a flat gate.

Flash memory An EPROM or EEPROM with the capability of block erasure of data in the memory array.

Flat zone The highly temperature-controlled region of a tube furnace.

Flip-chip joining A chip or package connection process where "bumps" of connecting metal are formed on the chip surface and the chip is "flipped" over for soldering to the package.

Foup (Front opening unified pod) A wafer carrier used in automated wafer fabrication lines. It is a mini environment and mates with process tools in order to

maintain wafer cleanliness.

Four-point probe A piece of electrical test equipment used to determine the sheet resistivity of a wafer.

Furnace A piece of equipment containing a resistance-heated element and a temperature controller. It is used to maintain a region of constant temperature with a controlled atmosphere for the processing of semiconductor devices.

Fuse A circuit component that can be blown to allow a desired memory cell or logic gate to be programmed.

Gallium arsenide (GaAs) The most common of compound semiconductor materials. It has the advantage of producing higher-speed devices than those produced using silicon as a substrate.

Gate Along with the source and drain, one of the three regions of the unipolar or field-effect transistor (FET). The gate controls the current flow between the source and drain.

Gate array Type of integrated circuit made up of an arrangement of interconnected gates used to provide custom functions.

Gate oxide (gate ox) The thin oxide that causes the induction of charge, creating a channel between source and drain regions of a MOS transistor.

Germanium Semiconducting material used in the manufacture of crystal diodes and of early transistors.

HEPA filter (high-efficiency particulate attenuator) A filter constructed of fragile fibers in an accordion-folded design that allows a larger filtering area at an air velocity low enough for operator comfort. This filter permits a filtering efficiency of 99.99 percent.

Hexamethyldisilazane (HMDS) Primer used to promote photoresist adhesion.

High-pressure oxidation Oxidation carried out at high pressure (10 to 20 atm) to reduce the amount of heat or time required. The reaction chamber for this process must be constructed of stainless steel to safely contain the pressure.

Hole (1) The absence of a valence electron in a semiconductor crystal. Motion of a hole is equivalent to motion of a positive charge. (2) A “hole” in a surface layer

created by the photomasking process.

Hybrid integrated circuit A structure consisting of an assembly of one or more semiconductor devices and a thin-film integrated circuit on a single substrate, usually of ceramic.

Hydrofluoric acid (HF) An acid used to etch silicon dioxide; often diluted or buffered before it is used.

Hydrogen (H₂) A gas used in semiconductor processing primarily as a carrier gas for high-temperature reaction steps such as epitaxial silicon growth.

Hydrophilic Affinity toward water (water-loving); a hydrophilic surface is one that will allow water to spread across it in large puddles.

Hydrophobic Aversion to water; a hydrophobic surface is one that will not support large pools of water. The water is pulled into droplets on the surface. These surfaces often are termed “dewetted.”

Hydroscopic Attracts and absorbs water.

Integrated circuit A circuit in which many elements are fabricated and interconnected on a single chip of semiconductor material—as opposed to a “nonintegrated” circuit, in which the transistors, diodes, resistors, and so on, are fabricated separately and then assembled.

Integration level The range of total component count in a die. Varies from SSI (small-scale integration, less than 50 components) to ULSI (ultra-large-scale integration, over 1,000,000 components).

International technology roadmap for semiconductors (ITRS) A roadmap of future requirements for wafer fabrication processes, factory operation, devices, materials, and functions.

Interposer A passive chip containing metallization and vias to allow the connection of separate die in a package.

Intrinsic semiconductor An element or compound that has four electrons in its outer ring (i.e., elements from Group IV of the periodic table or compounds of Group III and V).

Ion An atom that has either gained or lost electrons, making it a charged particle

(either negative or positive).

Ion beam milling A dry etching method that uses an ion beam. Argon atoms are ionized and accelerated toward a wafer. The exposed areas are removed through a sputtering action.

Ion implantation Introduction of selected impurities (dopants) by means of high-voltage ion bombardment to achieve desired electronic properties in defined areas.

Interconnect See [*lead*](#).

ISO 9000 The International Standards Organization standards for clean rooms.

Isolation diffusion Diffusion step resulting in P-N junctions surrounding the areas to be separated from each other.

Isotropic etching Refers to the etching of the photoresist both downward and to the side.

JFET (junction field-effect transistor) Device in which voltage is applied to a terminal to control current between the source and drain regions.

Junction The interface at which the conductivity type of a material changes from P type to N type or vice versa.

Killer defect A defect that causes the failure of a device or circuit.

Lateral diffusion The diffusion of dopants from side to side every time the wafer is heated near the diffusion temperature range.

Layering A process by which thin layers of different materials are grown on, or added to, the wafer surface.

Lead A metal strip on the wafer surface.

LED (light-emitting diode) A semiconductor device in which the energy of minority carriers in combining with holes is converted to light. Usually, but not necessarily, constructed as a P-N junction device.

Lift-off process A material removal process in which the material is deposited into a hole in a photoresist layer, and the pattern defined where the photoresist layer is removed (lifted off) the surface.

Light field mask See [clear field mask](#).

Lithography Process of pattern transfer; when light is utilized, it is termed photolithography; when patterns are small enough to be measured in microns, it is referred to as microlithography.

Local oxidation of silicon (LOCOS) A MOS surface isolation scheme in which silicon dioxide is grown around islands of silicon nitride, which is in turn removed to leave oxide-free areas for the formation of the circuit components.

Low-pressure CVD (LPCVD) A CVD system and process performed in a low-pressure range.

LSI (large-scale integration) Refers to chips with between 5000 and 100,000 components each.

Majority carrier The mobile charge carrier (hole or electron) that predominates in a semiconductor material—for example, electrons in an N-type region.

Mask A glass plate covered with an array of patterns used in the photomasking process. Each pattern consists of opaque and clear areas that respectively prevent or allow light through.

Masks are aligned with existing patterns on silicon wafers and used to expose photoresist. Mask patterns may be formed in emulsion, chrome, iron oxide, silicon, or a number of other opaque materials.

Masking See [patterning](#).

Memory The storing of data or information.

Metal mask The step at which an island of conducting material is left on the wafer surface.

Metalorganic CVD (MOCVD) A VPE process that uses halides and metalorganic sources.

Microchip See [chip](#).

Microelectromechanical systems (MEMS) Miniature (nanolevel) machines fabricated using semiconductor fabrication processes.

Micrometer One-millionth of a meter (10^{-6} m); abbreviation is μm .

Micron Same as micrometer.

Miller indices A numerical system of three numbers used to identify the orientation of planes in a crystal.

Mini-environment An environment that maintains wafer cleanliness by storing, transporting, and loading or unloading wafers in small, clean enclosures.

Minority carrier The nonpredominant mobile charge carrier in a semiconductor—for example, electrons in a P-type region.

MMOS (memory MOS) A nonvolatile memory device structure that enables information to be retained during power shutdown.

Mobile ionic contaminant Wafer contaminants that are electrically charged particles, which can cause electrical failures in a wafer or device.

Molded package A package of epoxy or other polymer material that is molded around the chip and package lead system.

Molecular beam epitaxy (MBE) An evaporative deposition process capable of extreme control of the deposition process.

Molecule Smallest quantity of a substance that retains the properties of that substance.

Monochromatic light Light of a single wavelength.

MOSFET A field-effect transistor containing a gate over thermal oxide over silicon.

MSI (medium-scale integration) Refers to chips with between 50 and 5000 components each.

Multichip module (MCM) A package containing two or more IC chips connected by a thin-film metal system.

Multilayer resist process An image-resolution process that uses two or more layers of photoresist.

Nanometer (nm) length unit = 1×10^{-9} meter.

Nanotechnology Processes and materials used to build semiconductor devices and other structures with nanometer dimensions.

Negative resist Photoresist that remains in areas that were not protected from exposure by the opaque regions of a mask while being removed by the developer in regions that were protected. A negative image of a mask remains following the develop process. A “clear” or “light” field mask is most often used with negative resist.

Next generation lithography (NGL) The processes, materials, and tools used to pattern wafers with feature sizes in the nanometer range.

Nitric acid (HNO_3) A strong acid often used to clean silicon wafers or etch materials.

Nitridation The formation of silicon nitride by the high-temperature exposure of a silicon surface to nitrogen.

Nitrogen (N_2) A gas that seldom reacts with other materials. It is often used as a carrier gas for chemicals in semiconductor processing.

NMOS N-channel MOS; type of MOSFET in which the channel is negative during conduction.

Nonvolatile memory circuit A memory circuit that retains its data when power to the chip is lost.

N-type A semiconductor material in which the majority of carriers are electrons and therefore negative. N-type dopants in silicon are Group V elements, in which the fifth outer electron is free to conduct current.

NPN transistor A transistor that has a base of P-type silicon sandwiched between an emitter and a collector of N-type silicon.

Ohm's law A relationship between resistance, voltage, and current; $R = V/I$.

Oil diffusion pump A type of high-vacuum pump that uses evaporated hot oil particles to “push” chamber particles out of the system.

Optical proximity masks Photo masks and reticle with patterns designed to account

for diffraction effects during the exposure process.

Overall yield The percentage of functioning packaged chips from a wafer related to the number of die mapped onto the wafer. Overall yield is the product of fabrication yield, sort yield, and assembly yields.

Oxidation The growth of oxide on silicon when exposed to oxygen. This process is highly temperature dependent.

Oxidation reaction chamber A chamber in which oxidation takes place. Quartz or silicon carbide tubing is used to make an oxidation reaction chamber due to their thermal resistance and purity.

Oxide See [*silicon dioxide*](#).

Oxide etching An etching process that uses acid—usually hydrofluoric acid (HF). The acid must be buffered for the reaction to proceed at a rate slow enough to be controlled. Buffered oxide etch (BOE) is often used.

Oxygen (O₂) Gas used to combine with silicon to form silicon dioxide.

Package Protective container for a semiconductor chip having electrical leads.

Packaging yield The percentage of packaged die passing the final tests as compared to the number of good die that entered the packaging process.

Passivation Sealing layer added at the end of the fabrication process to prevent deterioration of electronic properties through chemical action, corrosion, or handling during the packaging processes. The passivation layer, usually silicon dioxide or silicon nitride, protects against moisture or contamination.

Patterning A process in which the pattern in a reticle or photomask is transferred to a wafer, resulting in the identification of areas to be doped or selectively removed.

Pellicle A thin film of an optical-grade polymer that is stretched on a frame and secured to a mask or reticle. This solves the problem of airborne dirt collecting on the mask and acting as an opaque spot. During the exposure, any dirt is held out of the focal plane and does not “print” onto the wafer.

Phosphine (PH₃) A gas that is often used as a source of phosphorus for doping silicon.

Phosphorus (P) The N-type dopant commonly used for the sinker and emitter diffusions in standard bipolar integrated circuit technology.

Phosphorus oxychloride (POCl₃) A liquid that is often used as a source of phosphorus for doping silicon.

Photomasking See [patterning](#).

Photoplate A coated mask blank before imaging.

Photoresist The light-sensitive film spun onto wafers and “exposed” using high-intensity light through a mask. The exposed (or unexposed, depending on its polarity) photoresist is dissolved with developers, leaving a pattern of photoresist that allows etching to take place in some areas while preventing it in others.

Pin grid array (PGA) A large chip package with many leads coming out of the entire bottom surface of the package.

Pinhole A small, undesired hole in the photoresist or in the opaque region of a mask or reticle.

Planar structure A flat-surfaced device structure fabricated by diffusion and oxide masking, with the junctions terminating on a single plane.

Planarization Flattening of the wafer surface during the fabrication process by heat flow, organic layers, or chemical mechanical polishing techniques.

Plasma High-energy gas made up of ionized particles.

Plasma-enhanced CVD (PECVD) A CVD system and process using plasma energy to drive the deposition.

Plasma etch A dry-etch process using reactive gases energized by a plasma field.

Plastic package See [molded package](#).

Plug (via plug) A metal (generally a refractory metal) deposited in a connecting via hole between the conducting layers of a multilayer metal system.

PMOS (P-channel MOS) Type of MOSFET in which the channel is positive due to conduction achieved by holes.

PNP Semiconductor crystal structure consisting of an N-type region sandwiched between two P-type regions, as commonly used in bipolar transistors.

Polycide MOS gate A MOS gate structure composed of a sandwich of silicon dioxide topped with a layer of polysilicon that is covered with a refractory metal silicide.

Polycrystalline silicon (poly) Silicon composed of many crystal unit cells randomly arranged.

Polymer A complex organic chemical compound made up of repeating units.

Positive resist Photoresist that is removed in areas that were not protected from exposure by the opaque regions of a mask, while remaining after develop in regions that were protected from exposure. A positive image of the mask remains following the develop process. A “dark field” mask is used most often with positive resist.

Post-exposure bake (PEB) A baking step performed after resist exposure to reduce standing wave effects.

Predeposition (predep) The process step during which a controlled amount of a dopant is introduced into the crystal structure of a semiconductor.

Primer chemical A chemical that enhances the adhesion of a desired layer. In semiconductor technology, the layer is usually photoresist.

Process tool The term used for wafer fabrication process equipment and system.

Programmable read-only memory See [*PROM*](#).

Projection alignment An exposure system in which the image on the mask is projected onto the wafer. This results in little mask or photoresist damage and has about the same productivity as the contact method. For LSI and VLSI, projection alignment is standard.

Projection aligner An aligner tool that projects the mask or reticle image over a distance onto the wafer.

PROM (programmable read-only memory) A technology in which fuses are used in every memory cell, and selected fuses are blown so as to program the chip with user-specified information.

Proximity aligner An aligner tool that holds the wafer and mask or reticle a small distance apart during the resist exposure cycle.

P-type Semiconductor material in which the majority carriers are holes and therefore positive. P-type dopants in silicon are Group III-A elements.

Quartz Commercial name for silicon dioxide formed into glass products. Because of its high temperature resistance, quartz is used in many processing steps in integrated circuit fabrication.

RAM (random access memory) Device that temporarily stores digital information.

Rapid thermal oxidation (RTO) An oxidation process performed in a rapid thermal processing (RTP) tool (See below).

Rapid thermal processing (RTP) A single-wafer processing tool that uses high-intensity lights or other sources to heat and cool the wafer in milliseconds.

RCA clean A multiple-step process to clean wafers before oxidation; named after RCA, the company that developed the procedure.

Reactive ion etching (RIE) An etching process that combines plasma and ion beam removal of the surface layer. The etchant gas enters the reaction chamber and is ionized. The individual molecules accelerate to the wafer surface. At the surface, the top layer removal is achieved by the physical and chemical removal of the material.

Reactor (1) A piece of equipment used for the deposition of a layer of material used in semiconductor processing. Common types of reactors are epitaxial reactors, vapox reactors, and nitride reactors. (2) See [plasma etcher](#).

Refractory metals A group of metals with the common characteristics of resistance to heat, wear, and corrosion. They are used as conductors in plug/via systems and include molybdenum, tantalum, and tungsten.

Resistivity A measure of the resistance to current flow in a material. A function of the attraction between the outer electrons and inner protons of a material. The more tightly bound the electrons, the greater the resistivity.

Resolution capability The minimum feature size capability of a photolithography process or tool.

Reticle An exposure mask with only a portion of complete die pattern.

Rinse The removal of wet etchants or developers from the wafer. This process results in stopping the etching or developing processes and removing the active chemical from the surface. There are several different methods of rinsing: overflow rinsing, spray rinsing, dump rinsing, and spin-rinse dryers.

ROM (read-only memory) A memory circuit containing permanent data with no capability of accepting new information.

Salicide MOS gate A polycide MOS gate structure sequenced in the process to self-align to the source or drain. See [polycide MOS gate](#).

Scanning electron microscope (SEM) Microscope used to magnify images as much as 50,000 times by means of scanning with an electron beam. The impinging electrons cause electrons on the surface to be ejected. The ejected electrons are collected and translated into a picture of the surface.

Scribe lines Lines used to separate die on a wafer. The wafer will be sawed along the scribe lines, resulting in individual chips.

Self-aligned gate A MOS structure that allows the direct alignment of the source or drain to the gate without a photoresist alignment step.

Semiconductor An element such as silicon or germanium, intermediate in electrical conductivity between the conductors and the insulators, in which conduction takes place by means of holes and electrons. Common single-element semiconductors are Si (silicon) and Ge (germanium); a common compound semiconductor is GaAs (gallium arsenide).

Sheet resistance A measurement of resistance with dimensions of ohms per centimeter squared that shows the number of N-type or P-type donor atoms in a semiconductor.

Side diffusion see [lateral diffusion](#).

Silicon (Si) The Group IV element used for fabricating diodes, transistors, and integrated circuits.

Silicon dioxide (SiO₂) A nonconductive layer that can be thermally grown or deposited on silicon wafers. Thermal silicon dioxide is commonly grown using either oxygen or water vapor at temperatures above 900°C.

Silicon gate MOS An MOS gate structure with a layer of polysilicon on top of a

thin layer of silicon dioxide.

Silicon nitride (Si_3N_4) A nonconductive layer chemically deposited on wafers at temperatures between 600 and 900°C. When it is the final layer in the process, it protects devices against contamination.

Single crystal Refers to substances that have all unit cells arranged in a definite and repeated fashion as opposed to polycrystalline materials, which have unit cells randomly arranged.

Slope etching Controlled undercutting; an etch strategy in which the sides of the contact holes are purposely overetched so as to reduce the shadow effect of the sidewall and the resultant thinning of the film.

Soft baking A heating process used to evaporate a portion of the solvents in resist. The term “soft” describes the still-soft resist after baking. The solvents are evaporated to achieve two results: to avoid retention of the solvent in the resist film, and to increase the surface adhesion of the resist to the wafer.

Solid-state electronics Designation used to describe devices and circuits fabricated from solid materials such as semiconductors, ferrites, or thin films, as distinct from devices and circuits making use of electron tube technology.

Source Along with the gate and drain, one of the three regions of a unipolar or field-effect transistor (FET).

Spectrophotometer An analytical instrument used to collect interference measurements that are used to calculate film thickness.

Spinning A technique in which the photoresist is spun onto the wafer, resulting in a typical photoresist layer 0.5 μm thick with an allowable thickness variation of 10 percent.

Spin rinse dryer A machine that automatically rinses and dries wafers by spinning them in cassettes around a central axis.

Spray development A system of developing the pattern in a photoresist layer while the wafer is spinning on a vacuum chuck.

Spreading resistance A technique used for measuring the dopant concentration profile in a wafer.

Sputtering A method of depositing a thin film of material on wafer surfaces. A target of the desired material is bombarded with radio-frequency-excited ions that knock atoms from the target; the dislodged target material deposits on the wafer surface.

SSI (small-scale integration) Refers to chips with between 2 and 50 components each.

Standard mechanical interface (SMIF) A system that allows the mating of portable clean wafer boxes (called *pads*) to the clean microenvironment loading stations of process tools.

Standing wave effects A vertical resist exposure pattern that follows standing waves set up in the resist layer by constructive interference of the exposure light reflecting off the wafer surface.

Static RAM (static random access memory) Fast read-write memory cell based on transistors.

Steam oxide Thermal silicon dioxide grown by bubbling a gas (usually oxygen or nitrogen) through water at 98 to 100°C.

Step and repeat An operation in which the pattern on the reticle is transferred to the mask or wafer. The photoresist-coated mask blank (chrome, emulsion, or iron oxide) or wafer is placed on an x-y stage, and the reticle pattern is repeatedly imaged until the entire surface is filled with the reticle pattern.

Step coverage The ability of new layers to evenly cover steps formed in the existing wafer layers.

Stepper An aligner tool that aligns and exposes one (or a small number) of die at a time. The tool “steps” to each subsequent die on the wafer.

Stripping Removal process; usually refers to photoresist.

Subcollector See [*buried layer*](#).

Substrate The underlying material upon which a device, circuit, or epitaxial layer is fabricated.

Sulfuric acid (H₂SO₄) A strong acid often used to clean silicon wafers and to remove photoresist.

Susceptor The flat slab of material (usually graphite) upon which wafers are held during high-temperature deposition processes such as epitaxial growth or nitride deposition.

System in package (SIP) A collection of chips in a single package that includes the functions of an electronic system.

System on chip (SOC) A single chip with different sections (logic, memory, etc.) that function as a complete electronic system.

Tape automated bonding (TAB) A chip-to-package connection process in which the package leads are formed on a flexible tape, and all the lead fingers are bonded to the chip in one action.

Target The material to be sputtered during the sputtering process.

TCE (trichloroethylene) A solvent used for wafer and general cleaning.

Test die Die on a wafer that appear to have a different pattern from most others. These contain test devices created by the same processes at the regular die; however, the devices on these die are designed on a larger scale to allow in-process quality control.

Tetraethylorthosilicate (TEOS) A chemical source for the deposition of silicon dioxide.

Thermal diffusion A process by which dopant atoms diffuse into the wafer surface by heating the wafer in the range of 1000°C and exposing it to vapors containing the desired dopant.

Thermal oxide On silicon semiconductor devices, an oxide fabricated by exposing the silicon to oxygen at high temperatures. The resulting interface has low levels of ionic impurities and defects (surface states).

Thermocouple A device to measure the temperature in a furnace of a reactor. It is made by welding two wires together at a point. Heat generates a voltage between the two materials that is proportional to the temperature.

Through silicon via (TSV) A hole created and filled (plug/via) through a silicon chip to facilitate a top-to-bottom metal connection system.

Torr Pressure unit; international standard unit replacing the English measure,

millimeters of mercury (mmHg).

Transistor A semiconductor device that uses a stream of charge carriers to produce active electronic effects. The name was coined from the electrical characteristic of “transfer resistance.”

Tube (1) See [furnace](#). (2) A cylindrical piece of quartz with fittings on one or both ends. It is placed in a furnace to provide a contamination-free and controlled atmosphere.

ULSI (ultra-large-scale integration) Refers to chips with more than 1,000,000 components each.

Ultraviolet (UV) light A portion of the electromagnetic spectrum from 250 to 500 nm. High-pressure mercury sources emit UV light for photoresist exposure.

Undercutting See [isotropic etching](#).

Unipolar transistor A transistor such as an FET whose action depends on majority carriers only.

Vacancy (1) A position in the crystal for an atom which is empty. (2) A type of crystal defect.

Vacuum A low-pressure condition.

Vapor phase epitaxy (VPE) An epitaxial deposition system that can combine several source gases to deposit compound semiconductors.

Vapor priming A technique in which primer is applied in a vapor state such that the wafer never comes in contact with any possible contamination in the liquid—or in the case of HMDS, with any particles of hydrolyzed HMDS.

Vertical tube furnace An oxidation, diffusion, or other tube process with the tube oriented in a vertical position. These systems provide increased temperature zones and smaller foot prints.

Via Vertical opening filled with conducting material used to connect circuits on various layers of a device to one another and to the semiconducting substrate; serves same purpose as “contacts.”

Viscosity The qualitative measure of liquid flow. Viscosity measurements are made

by measuring the force required to move an object through the liquid. It is a measurement of “internal friction.”

VLF hood A workstation with vertical laminar airflow to keep particulate levels low.

VLSI (very-large-scale integration) Refers to chips with between 100,000 and 1,000,000 components.

Volatile memory circuit A memory circuit that loses its data when power to the chip is lost.

Voltage The force applied between two points causing charged particles (and hence current) to flow.

Wafer A thin, usually round, slice of a semiconductor material from which chips are made.

Wafer fabrication The series of manufacturing operations in which the circuit or device is put in and on the wafer.

Wafer flat Flat area(s) ground onto the wafer’s edges to indicate the crystal orientation and the dopant type of the wafer structure.

Wafer sort The step after wafer fabrication during which the electrical parameters of integrated circuits are tested for functionality. Probes contact the pads of the circuit to conduct the test, leading to the name “prober” for the equipment that performs electrical tests on each die site of completed wafers.

Wafer sort yield The number of functioning die at wafer sort as compared to the total number of die started; typically, the lowest major yield point for integrated circuits.

Wire bonding An assembly step in which thin gold or aluminum wires are attached between the die bonding pads and the lead connections in the package.

X-ray aligner An aligner tool that uses X-rays and a mask to expose resist-coated wafers.

X-ray exposure system Imaging system using X-rays as the exposure source. Due to their short wavelengths, X-rays exhibit no detrimental diffraction effects.

Yield A percentage used in the semiconductor industry that indicates the amount of finished product leaving a process as compared to the amount of product entering that process.

A

Acceleration

Acceleration tube, ion implantation system

Acceptors

Accumulative wafer-fabrication yield

Acids

Acoustic streaming

Acoustic wave devices

Acousto-optical modulator (AOM)

Additives, photoresist

Adhesion capability, photoresist

Adhesive floor mats

Adiabatic cooling

Advanced lithography

Aerial image

Aerosols

AFM (atomic force microscopy)

AGVs (automated guided vehicles)

Air, contaminants in

Air pressure, cleanroom construction and

Air showers

Airborne molecular contaminants (AMCs)

Aligners:

contact

projection

types of

Alignment

of die to package

and exposure (A&E)

Alignment criteria

Alignment errors (misalignment)

Alignment marks (targets)

Alkaline-water solutions

Alkalis (bases)

Alloying

Alternating phase-shift mask (alternating aperture phase shift mask-AAPSM)

Aluminum:

as conductor

voltage/current (V/I) versus thickness of

Aluminum wire bonding
Aluminum wiring
Aluminum-copper alloy
Aluminum-film wet etching
Aluminum-silicon alloys
Amorphous materials
Amorphous silicon deposition, polysilicon and
Amplification
Amplified resist
Amplifier circuits
Analog logic circuits
Analyzing, mass, ion implantation system
Anisotropic etching
Annealing
 and dopant activation
Annular-ring illumination
Antireflective coatings (ARCs)
APCVD (atmospheric-pressure CVD) systems
ARCs (antireflective coatings)
Argon:
 cryogenic cleaning
 cryokinetic cleaning
 ion-beam etching
 sputtering
Arrays
Arsenic, as dopant
Arsenic buried layers
Ashing
Aspect ratio, photoresist
Assay number
Assembly(See Packaging)
Assembly and test (A/T)
Atmospheric-pressure CVD systems (APCVD)
Atomic force microscopy (AFM)
Atomic layer deposition (ALD)
Atomic number
Atomic structure
Attachment to package (or substrate)
Auger electron spectroscopy
Auger electrons
Autodoping, P-type film
Automated guided vehicles (AGVs)
Automatic defect detection
Automatic spinners
Automatic wafer loading

Automation

- closed-loop control-system**
- factory-level**
- process**
- wafer-loading**
- (See also Wafer fabrication)**

Axial dryer

B

Back end of the line (BEOL)

- cleaning**

Back-end process (see Packaging)

Backside coating

Backside damage

Backside gold

Backside metallization

Backside processing

Backstreaming

Bacteria, as contaminants

Bake (baking):

- hard**
- inside-out**
- microwave**
- post-exposure**
- post-soft-bake cooling (chill)**
- soft**
- vacuum**

Ball bonding

Ball flip-chip bonding

Ball-grid arrays

Ballroom design

Barrel radiant-induction-heated APCVD

Barrel-radiant-heated PECVD

Barrier deposition

Barrier layer

Bases (alkalis)

Batch versus single-wafer processing

Batch-immersion etching

Batteries, thin film rechargeable

Bays, service

BCDSs (bulk chemical distribution systems)

Beam focus

Beam scanning

Bell Labs

BEOL (back end of the line)
Beta
Beveling
Biasing
Bi-MOS (or bi-CMOS) circuits
Binary notation
Bipolar circuits
Bipolar transistors
Bird's beak
Bits
Blow-off guns
Body covers
Bohr atom model
Bonded wafers
Bonding processes

- bump or ball flip-chip bonding**
- die attach**
- tape automated bonding (TAB)**
- wire bonding**

Book-to-bill ratio (b/b)
Boron

- P-type dopants**

Boron silicate glass (BSG)
Boron slugs
Borophosphorus silicate glass (BPSG)
Breakdown voltage
Breath, contamination control and
Bridges between adjacent patterns
Brown pads
Bubblers
Budgets
Buffered oxide etches (BOEs)
Bulk chemical distribution systems (BCDSs)
Bump or ball flip-chip bonding
Bunny suits (oversuits)
Buried layers
Burn-in tests

C

CAD (computer-aided design)
Calcium fluoride
CAM (computer-aided manufacturing)
Capacitance-voltage (C/V) analysis
Capacitors

- junction**
- metallization**
- MOS transistors**
- oxide-silicon**
- stacked**
- trench**

Capillary

Capillary force

CAPP (computer-aided process planning)

Caps

Carbon bed filtration

Carrier illumination junction depth

Carrier mobility

Carriers

Carro's acid

Cascade rinsers

Cavitation

CEL (contrast enhancement layer)

Cell plate

Cellulose acetate (AC) films

Celsius or centigrade scale

Centerless grinder

Centipoise

Centistoke

CERDIP (ceramic DIP) packages (Cerpacks or Cerflats)

Channel lengths

Channel stops

Channeling

Chelating agent

Chemical cleaning

Chemical vapor deposition (see CVD (chemical vapor deposition) systems)

Chemical-cleaning solutions

Chemically amplified resists

Chemical-mechanical polishing (CMP)

- basic processing steps**

- dual-damascene process**

- planarity and**

- polishing pads**

- post-CMP clean**

- slurry**

- summary**

- tools**

Chemical-mechanical processing

Chemicals:

- as contaminants**

- process
- purity of
- spray cleaning

Chill (post-soft-bake cooling)

Chip sizes, increasing

Chip-scale packages

Chlorine-added oxidation

CIM (computer-integrated manufacturing)

Circuit components (see Integrated circuits)

Circuit design

(See also Integrated circuits)

Circuit layout

Class numbers of areas

Clean air strategies

Cleaning:

- chemical
- chemical-cleaning solutions
- cryogenic
- deflux
- dry
- high-pressure water
- post-CMP clean
- preoxidation wafer cleaning
- RCA cleans
- room temperature and ozonated chemistries
- sonic-assisted
- spray
- wafer-surface

(See also Contaminants; Rinsing)

Cleanliness

- oxide and furnace
- packaging and
- of photoresists

Cleanrooms

- construction of
 - construction materials
 - glove cleaners
 - shoe covers
 - static control
- maintenance of
- materials and supplies
- personnel-generated contamination
- traditional

(See also Contaminants)

Clear-field masks

- Closed-loop control-system automation
- Clothing, contamination control and
- Cluster arrangement
- Clustering
- CMOS (complementary metal oxide silicon) circuits
 - epitaxial films
 - memory circuits
- Coating, backside
- Cold-wall systems
- Collector current
- Collimated beam (collimator; collimated light)
- Collimated light inspection
- Color, layer thickness measurements and
- Color versus thickness chart
- Columnar poly
- Components (see Integrated circuits; Semiconductor devices)
- Composite drawing
- Compounds
 - semiconducting
- Computer-aided design (CAD)
- Computer-aided manufacturing (CAM)
- Computer-aided process planning (CAPP)
- Computer-integrated manufacturing (CIM)
- Concentration-versus-depth graphs
- Conduction
- Conductivity
 - type of (N or P)
- Conductors
 - materials for
 - underpass
- Confocal microscopes
- Conformal layer
- Constant acceleration
- Construction materials, cleanroom
- Contact aligners
- Contact holes (contacts)
- Contactless C/V measurement
- Contaminants (contamination control)
 - air
 - airborne molecular contaminants (AMCs)
 - bacteria
 - chemicals
 - clean air strategies
 - cleanroom workstation strategy
 - detection of

- drying techniques
- equipment
- etching
- metallic ions
- micro and mini environments
- particles
- particulate removal
- personnel-generated
- photoresists
- problems caused by contaminants
- process chemicals
- process water
- quartz
- sources of
- temperature, humidity, and smog
- tunnel or bay design concept
- types of contaminants
- wafer scrubbers
- wafer-surface cleaning
- water rinsing
- (*See also* Cleaning; Cleanrooms; Rinsing)

Continuous conduction-heated APCVD system

Contrast effects

Contrast enhancement layer (CEL)

Contrast threshold

Convection

Convection ovens, for soft baking

Cooling, post-soft-bake (chill)

Copper, contamination control

Copper metallization (damascene) bump bonding

Copper wiring

Copper-dual damascene process

Corrosion, plasma etching

Cost of ownership (CoO)

- plasma etching

Costs:

- of chips

- labor

- wafer fabrication

Critical dimension (CD) budgets

Critical dimensions

Crusting:

- effect of ovens

- resist stripping

Cryogenic cleaning

- Cryokinetic cleaning
- Crystal damage, ion implantation and
- Crystal defects (dislocations)
- Crystal flat grinding
- Crystal growth
- Crystal impurities
- Crystal orientation
- Crystal planes
- Crystals
 - poly and single
 - quality
- C/V (capacitance-voltage) analysis
 - contactless C/V measurement
- CVD (chemical vapor deposition) systems
 - atmospheric-pressure systems
 - atomic layer deposition
 - basic components of
 - deposited films
 - high-density plasma (HDPCVD)
 - layer deposition
 - low-pressure (LPCVD)
 - metallization
 - metalorganic (MOCVD)
 - process steps
 - refractory deposition
 - types of
 - ultrahigh vacuum (UHV-CVD)
- Czochralski (CZ) method

D

- Dark space
- Dark-field inspection
- Dark-field masks
- Dash necking
- Deep ultraviolet (DUV)
- Defect density:
 - die area and
 - reduction in
 - yields and
- Defects:
 - automated inline inspection systems
 - crystal
 - detection of
 - growth

- killer
- mask
- point
- process
- yields and sizes of

Deflashing

Deflux

Deforest, Lee

Dehydration baking

Deionized (DI) water

Demand-pull system

Densification

Density (of chips)

- defined

- vapor

Dep (see Deposition)

Depletion

Deposited-oxide wet etching

Deposition

- defined

- (*See also* Layer deposition)

Depth, measurement of

Depth of focus (depth of field)

- scanning electron microscopes and

Descum, plasma

DESIRE process

Develop inspect (DI)

- causes for rejecting wafers

- manual inspection

- methods

- reject categories

Development

- dry (or plasma)

- negative resist

- positive resist

- puddle

- spray

- wet

Device electrical measurements

Device performance, contamination and

Device processing yield

Device reliability, contamination and

Devitrification

Diameter grinding

Diameter of wafers

- crystal defects and
- die sizes and
- edge die and
- increasing
- process variations and
- yields and

Diamond scribing

Diamond semiconductors

Diatomic molecules

Dichlorosilane, source chemistry

Die area, defect density and

Die separation

Die shrink (scaling)

Die sizes, wafer diameter and

Die sort (*see* Wafer sort)

Die-attachment area

Dielectric constant

Dielectric isolation

Dielectric strength

Dielectric wearout

Dielectrics

- high-k and low-k

- silicon dioxide

Die-on-die technique

Die-on-wafer technique

Dies (dice)

- alignment of die to package

- attachment

- attachment to package (or substrate)

- good die problem

- inspection of

- pick and place operation

- stacking die techniques

Die-to-database system

Differential oxidation rates

Diffraction

Diffused junctions

Diffused MOS (DMOS)

Diffusion doping

- concept of

- deposition

- formation of a doped region and junction

- lateral diffusion

- solid-state

Diffusion limited reaction

- Diffusivity
- Digital circuits
- Dimensional control
- Dimethylacetamide (DMAC)
- Dimethylformamide (DMF)
- Dimethylsulfoxide (DMSO)
- Diodes
 - doped, 422–423
 - Schottky barrier
- DIPs (dual in-line packages)
- Direct displacement vapor dryer
- Direct writing
- Discrete devices
- Dislocations
 - high-pressure oxidation and
- Dislocations (crystal defects)
- Dissolution inhibitor systems
- Donors
- Dopant concentration profile
- Doped diodes
- Doped junction isolation
- Doped polysilicon
- Doped region, formation of
- Doped resistors
- Doped semiconductors
- Doping (dopants)
 - concentration versus depth graphs
 - diffusion (*see* Diffusion doping)
 - epitaxial films
 - graphical representation of junctions
 - ion implantation (*see* Ion implantation)
 - N-type
 - oxidation effects
 - process goals
 - P-type
 - redistribution of
 - same-type doping
 - silicon dioxide and
 - sources of
- Dose, ion implantation
- Dosimetry, optical
- Double gates (DGs)
- Double masking technique
- Double-door pass-throughs
- Double-sided polishing

- Downstream plasma processing
- Downstream strippers
- Drag-type pumps
- DRAMs (dynamic random-access memories)
- Drawback (suckback)
- Drift
- Drive-in oxidation
- Dry cleaning
- Dry (plasma) development
- Dry etching
 - planar plasma etching
 - plasma etching
 - resist effects in
- Dry mechanical pumps
- Dry oxidation (dryox)
- Dry oxygen
- Dry stripping
- Drying techniques
- Dual in-line packages (DIPs)
- Dual-damascene copper process
- Dual-damascene process
- Dump rinsing
- Dyed resists
- Dynamic dispense
- Dynamic random access memory (DRAMs)
- Dynamic spin dispensing technique

E

- E-beams (electron beams)
- Edge bead,181
 - removal of
- Edge chips
- Edge grinding and polishing
- EEPROM (electronically erasable PROM)
- Effusion cells
- Electrical conduction:
 - conductors
 - electron and holes and
- Electrical sort (*see* Wafer sort)
- Electrical testing
- Electrochemical plating (electroplating)
- Electrodes:
 - capacitor
 - RF

- Electrolytic plating
- Electrolytic staining
- Electromigration
- Electron beam aligners
- Electron beams (e-beams)
- Electron flood gun
- Electron spectroscope for chemical analysis (ESCA)
- Electron-beam projection lithography (EPL)
- Electronics industry:
 - organization of
 - (*See also* Semiconductor industry)
- Electrons
 - Auger
 - holes and
- Electrostatic discharge (ESD)
- Elements
 - periodic table of
- Ellipsometers
- Emission microscopy
- Emissivity
- Encapsulation
- Enclosures
 - (*See also* Packaging)
- End cropping
- End station
- Energy-sensitive polymers
- Engineered substrates
- Engineered wafers (substrates)
- ENIAC (Electronic Numeric Integrator and Calculator)
- Environmental, safety, and health (ESH) programs
- Environmental protection of chips
- Environmental tests
- EPI resistors
- Epitaxial films (or layers)
 - for CMOS circuit wafers
 - doping
 - ion implantation and
 - metalorganic CVD (MOCVD) and
 - process
 - quality
 - selective growth of
- Epitaxial silicon
- Epoxy die attach
- Epoxy molding
- EPROM (erasable programmable ROM)

- Equipment, contamination control and
- Equipment standards
- Error budget
- Error function
- ESD (electro-static discharge)
- Etch definition, improving
- Etch profile control
- Etch rate of plasma systems
- Etchback planarization
- Etching
 - batch-immersion
 - dry
 - goals and issues
 - incomplete etch
 - ion-beam
 - overetch and undercutting
 - plasma
 - selectivity
 - vapor
 - wet
 - wet-spray
- Etch-resistant photoresists
- Eutectic die attach
- Excimer lasers
- Exponential model
- Exposure
 - alignment and (A&E)
- Exposure sources
 - next-generation lithography
- Exposure speed, sensitivity, and exposure
 - source of photoresists
- Extreme ultraviolet (EUV) exposure source

F

- Fab floor layout
- Fabs (*see* Wafer fabrication)
- Factory-level automation
- Fahrenheit scale
- Fairchild Camera
- Faraday cup
- Feature size
 - decreasing
- Feature sizes, yields and
- Federal Standard 209E

FEOL (front end of the line)
Ferroelectric materials
Ferroelectric memories
FET (field effect transistor)
Field effect transistor (FET)
Field oxide (FOX)
Field oxides
Field programmable gate arrays (FPGAs)
Field-effect transistors
 junction (JFETs)
Final inspection
Final tests
FinFET
FinFet
First-fail basis
Flash memories
Flat packs (FPs)
Flatness (flatting process)
Flats, major and minor
Flip-chip ball-grid arrays
Flip-chip bonding
 bump or ball
Float-zone crystal growth
Flood guns
Floor mats, adhesive
Fluorescence microscopes
Focused ion beams (FIB)
Forward voltage
FOUP (front-opening unified pod)
FOUPs (Front Opening Universal Pods)
450-mm wafers
Four-point probe
FPGAs (field programmable gate arrays)
Front end of the line (FEOL)
Front Opening Universal Pods (FOUPs)
Full custom logic circuits
Fuller, Buckminster
Functional test

G

Gallium arsenide (GaAs)
 on silicon
 wafer breakage
Gallium-arsenic-phosphide (GaAsP) wafers

Gas control panel (gas-flow controller)
Gas source MBE (GSMBE)
Gases
Gas-phase cleaning
Gate arrays
Gate oxide integrity (GOI)
Gate stacks
Gate widths
Gates
Gaussian distribution
Germanium
Gettering
Glass damage
Global planarization
Glove cleaners
Gloves
GOI (gate oxide integrity)
Gold wire bonding
Good die problem
Gowning, order of
Gowning area
Gowns
Green fabs
Grinding:
 edge
 orientation indicators
Groove (or bevel) and stain technique
Gross leak testing
Grove, Andrew
Grove, Sir William Robert
Growth defects
Growth rate, silicon dioxide
GSMBE (gas source MBE)

H

Half pitch
Halides, metalorganic CVD and
Handler
Hard bake
Haze
HDA (hydroxylamine)
Headgear
Heads, of manual spinners
Heat dissipation, packaging and

Heat sensitivity, of photoresists
Heat treatments
Heat-transfer methods, soft bake and
HEPA (high-efficiency particulate attenuation) filters
Hermetic sealing
Heteroepitaxial films
Heterojunctions
Heterostructures
Hexamethyldisilazane (HMDS)
HF-last surfaces
High-aspect-ratio patterns
High-density plasma CVD (HDPCVD)
High-k dielectrics
High-pressure mercury lamp sources
High-pressure oxidation
High-pressure water cleaning
Hi-rel packages, processes, and tests
Histograms
Hi-vac pumps
HMDS (hexamethyldisilazane)
Hole flow
Holes, electrons and
Homoepitaxial films
Hoods
Horizontal conduction-convection-heated LPCVD system
Horizontal conduction-heated APCVD system
Horizontal tube furnaces
Horizontal vertical-flow PECVD
Horizontal-tube induction-heated APCVD
Horni, Jean
Hot lots
Hot plates:
 in-line, single-wafer
 manual
 moving-belt
Hot-wall systems
Humidity, contamination control and
Hybrid circuits
Hydrated surfaces
Hydrofluoric acid (HF)
Hydrogen ions
Hydrogen peroxide:
 sulfuric acid with
 wet chemical stripping
Hydrogen reduction of trichlorosilane

Hydrophilic surfaces
Hydrophobic surfaces
Hydroscopic surfaces
Hydroxylamine (HDA)

I

IFET (insulated field effect transistor)

Illumination:

annular-ring
off-axis

Image reversal

Immersion cleaning

Immersion development method

Immersion exposure system

Implanted regions, dopant concentration in

Impurities

oxide
(*See also* Contaminants)

Incomplete etch

Index of refraction

Induction

Infrared ovens, moving-belt

Ingots

Inner leads

Inorganic acids

Inorganic residues

Inside diameter (ID) saws

Inside-out baking

Inspection:

develop inspect (*see* Develop inspect)
final
manual
postbonding and preseal
(*See also* Tests and measurements)

Insulators

silicon dioxide

Integrated circuits (ICs)

bi-MOS (or bi-CMOS)
CMOS (complementary MOS)
ferroelectric memories
hybrid circuits

introduction to

linear

- localized oxidation of silicon (LOCOS)

- logic circuits

- memory circuits

- MOS

- next generation of

- nonvolatile memories

- RAM (random-access memory)

- redundancy

- shallow trench isolation

- silicon on insulator (SOI)

- standard

- system-on-a-chip (SOC)

- terminology

- types of

- volatile memories

- (*See also* Semiconductor devices; Integrated circuits)

Integrated Device Manufacturers (IDMs)

Integrated processing

Integration level

Intel Corporation

Interconnection levels, increase in

Interconnects

Intermediate metal dielectric (IMD)

Intermetallic dielectric (interdielectric) layers (IDL or IMD)

International Technology Road map for Semiconductors (ITRS)

- maximum defect densities

- “nodes” of future devices

- Yield Enhancement

Interposers

Intrinsic semiconductors

Inventory control

Inversion voltage (threshold voltage)

Inverted surface

Ion densities, plasma system

Ion implantation

- acceleration tube

- analog circuits and

- annealing and dopant activation

- beam focus

- beam scanning
- channeling
- concept of
- crystal damage
- dopant concentration in implanted regions
- drawbacks of
- end station and target chamber
- evaluation of implanted layers
- implant specie sources
- ionization chamber
- masks
- mass analyzing or ion selection
- neutral beam trap
- plasma ion immersion
- system for
- uses of
- wafer charging

Ion milling

Ion plantation

Ion separation

Ion-beam etching

Ionization chamber

Ionized deposition (I-PVD)

Ionizers

Ions

ISO 9000 guidelines/standards

ISO Global Cleanroom Standards (ISO 14644-2)

Isolation:

- dielectric
- doped junction
- junction
- MOS LOCOS
- trench
- wafer

Isolation scheme, ion implantation

Isopropyl alcohol vapor drying

Isotropic etching

J

JFETs (junction field-effect transistors)

JIT (just-in-time inventory control)

Josephson, B. D.

Josephson junction

Junction capacitors

Junction delineation
Junction depth measurement
Junction field-effect transistors (JFETs)
Junction isolation
Junctions:

- defined
- formation of a doped region and
- graphical representation of
- Josephson
- N-P
- ultra-shallow

Just-in-time inventory control (JIT)

K

Kelvin scale
Kern, Werner
Kilby, Jack
Kilby circuit
Killer defects
Kinematic viscosity

L

Labor costs
Laminar gas flow
Laser dots
Laser signal control
Laser-assisted discharge plasma (LDP)
Lasers, excimer
Latch-up
Latent image
Lateral diffusion
Lattice
Layer deposition

- arsenic buried layers
- conductors
- film parameters
- gallium arsenide on silicon
- molecular beam epitaxy
- SOS (silicon on sapphire) and SOI (silicon on insulator)
- vapor-phase epitaxy (VPE)
- (*See also* CVD (chemical vapor deposition) systems; Epitaxial films)

Layer thickness measurements
Layering

LDD (lightly doped drain extension)

Lead on chip (LOC) package

Lead plating

Leads, package

electrolytic plating

inner and outer

plating process flows

trimming

Leakage current

Leaks, package

Lecture bottles

LEDs (light-emitting diodes)

Lenses:

numerical aperture of

reflection systems and

variable numerical aperture

LER (line edge roughness)

Less hermetic packages

Lift-off process

Light scattering, resist

Light sensitivity, of photoresists

Light-emitting diodes (LEDs)

Light-field masks

Lightly doped drain extension (LDD)

Light-sensitive polymers

Line edge roughness (LER)

Line organization, wafer fabrication

Linear circuits

Linear growth rate of silicon dioxide

Line-of-sight heat transfer method

Liner deposition

Line-width measurements

Liquid-encapsulated Czochralski (LEC)

crystal growing

Liquids

Lithography:

advanced

electron-beam

next-generation (*see* Next-generation lithography)

LOCOS (localized oxidation of silicon)

high-pressure oxidation and

Logic circuits

full custom

Logic diagram

Logic gates

Low-k dielectrics

defects, sources of
materials

LPCVD (low-pressure chemical vapor deposition) systems

CVD refractory deposition
polysilicon layers

M

Magnetron sputtering

Magnification, scanning electron microscopes

Major flats

Manometer

Manual hot plates

Manual inspection

Manual spinners

Manufacturing, stages of

Marangoni drying

Marking packages

Masks:

defects
ion implantation
phase-shift
(*See also* Patterning (photomasking))

Mass analyzing, ion implantation system

Mass interference

Mass-flow meter

Material preparation

Material Safety Data Sheet (MSDS)

Materials, wafer fabrication

Matter

properties of
states of

Maximum solid solubility

MBE (molecular beam epitaxy)

MCM (multichip module)

Mean free path

Measurements (*see* Tests and measurements)

Mechanical test wafers

Megachips

Megasonic cleaning

Melt

Memories (memory circuits)

ferroelectric
flash

nonvolatile

Memory MOS (MMOS)

MEMS (microelectromechanical systems)

MESFETs (metal semiconductor field-effect transistors)

Metal cans

Metal lines

Metal-gate MOSFET transistors

Metallic ions:

as contaminants

plasma resist stripping and

Metallization

backside

barrier metals

barrier or liner deposition

chemical-mechanical processing

conductors materials

copper-dual damascene process

CVD metal deposition

CVD refractory deposition

deposition methods

dry mechanical pumps

electrochemical plating

low-k dielectric materials

MOS gate and capacitor electrodes

multilevel metal schemes

plugs

refractory metals and refractory metal silicides

seed deposition

single-layer metal systems

sputtering

turbomolecular hi-vac pumps

vacuum systems

Metallized surfaces, wet chemical stripping of

Metallurgical microscope

Metalorganic CVD (MOCVD)

Metal-oxide-semiconductor field effect transistor (*see* MOSFET)

Metrology

Microelectromechanical systems (MEMS)

Micro-environments

Microloading

Microprocessors

Microscopes:

atomic force (AFM)

confocal

emission

- fluorescence
- metallurgical
- optical profilometry
- phase contrast
- scanning capacitance
- scanning electron (SEM)
- transmission electron (TEM)

Microwave baking

MICs (*see* Mobile ionic contaminants)

Mil-Standard 883

Mini-environments

Mini-fabs

Minor flats

Misalignment (alignment errors)

Mix and match aligners

Mixtures

MMOS (memory MOS)

Mobile ionic contaminants (MICs)

- capacitance-voltage plotting
- C/V evaluation

MOCVD (metalorganic CVD)

Molded epoxy enclosures

Molecular beam epitaxy (MBE)

Molecules

Momentum, defined

Momentum transfer

Monitor wafers (test wafers)

Monolithic Memories

Monolithic technology

Moore, Gordon

Moore's law

- wafer-fabrication business and

MOS (metal oxide semiconductor) devices (transistors)

- diffused (DMOS)
- ion implantation and
- measurement of
- memory (MMOS)
- N-channel
- P-channel
- polycide-gate
- salicide-gate
- silicon-gate
- sodium contamination

MOS gate oxides

MOS gates, metallization and

MOS LOCOS isolation
MOSFET (metal-oxide-semiconductor field effect transistor)
 alternatives to scaling challenges
 measurements
 metal-gate
Moving-arm dispensing
Moving-belt hot plates
Moving-belt infrared ovens
Muffle
Multichip modules
Multilayer resist
Multilevel metal schemes
Multivariable experiment analysis
Murphy, B. T.

N

Nanoanalysis era
Nanotechnology
National Semiconductor
N-channel transistors
Negative acting resists
Negative binomial model
Negative resist development
Negative resists, comparison of positive resists and
Neutral beam trap
Neutrons
Next-generation lithography (NGL)
 antireflective coatings (ARCs)
 challenges of
 chemical-mechanical polishing
 (see Chemical-mechanical polishing)
 contrast effects
 contrast enhancement layers
 dyed resists
 electron beam or direct writing
 etch profile control
 excimer lasers
 exposure issues
 exposure sources
 extreme ultraviolet (EUV)
 high-pressure mercury lamp sources
 image reversal
 improving etch definition
 lens issues and reflection systems

- numerical aperture of a lens
- optical process correction
- pellicles
- photoresist process advances
- reflow
- resolution challenges and solutions
- surface problems

Nitric acid, wet chemical stripping

Nitridation

Nitrocellulose (NC) films

Nitrogen bubbles

N-methyl pyrrolidine (NMP)

No-junction devices

Nonhermetically sealed packages

Non-HF-last process

Nonmetallic surfaces, wet chemical stripping of

Nonvolatile memories

Normal curve

Noyce, Robert

N-P junctions

NPN transistors

N-type dopants

N-type semiconductors

Nucleation

Number of process steps

Numerical aperture of a lens

O

Occupational, Safety, and Health Administration (OSHA)

OCD (optical critical dimension)

Off-axis alignment

Off-axis illumination

Ohm-centimeters (W-cm)

Ohmic contact

Ohms

Ohms (W)

Ohm's law

OHV (overhead hoist vehicle)

Open-loop lamp control

Optical aligners

Optical critical dimension (OCD)

Optical dosimetry

Optical image-shearing dimension measurement

Optical profileometers

Optical Proximity Corrected (or Optical Process Correction; OPC)

Optical resists

Optically modulated optical reflection (thermawave)

Optoelectronics

Order, color

Organic acids

Organic residues

OSHA (Occupational, Safety, and Health Administration)

Out of phase

Out-diffusion

Outer leads

Overall process yields

Overetching

plasma etching

Overflow rinsers

Overhead costs, wafer fabrication

Overhead hoist vehicle (OHV)

Overlay budget

Overshoot

Oversuits (bunny suits)

Oxidants, sources of

Oxidation

chlorine-added

dopants and

drive-in

dry (dryox)

of polysilicon

surface passivation and

(See also Silicon dioxide; Thermal oxidation)

Oxidation rates, differential

Oxide impurities

Oxide layer removal

Oxide masking

Oxide thickness

Oxide-nitride-oxide sandwiches (ONO)

Oxide-silicon capacitors

Oxygen, dry

Oxygen passivation (OP)

Ozone

cleaning with

P

Package leads (*see* Leads)

Package on package (package in package)

Packaging

- aluminum wire bonding**
- ball-grid arrays or flip-chip ball-grid arrays**
- bonding processes**
- bump or ball flip-chip bonding**
- CERDIP packages**
- chip characteristics and**
- chip-scale packages**
- cleanliness and static control**
- common package parts**
- deflashing**
- design of packages**
- die attach**
- die inspection**
- die pick and place**
- die separation**
- electrical testing**
- encapsulation**
- enclosures**
- environmental tests**
- epoxy die attach**
- eutectic die attach**
- final testing**
- functions and design**
- gold wire bonding**
- good die problem**
- heat dissipation**
- hybrid circuits**
- lead on chip (LOC) packages**
- leads**
- marking packages**
- metal can**
- metal cans**
- molded epoxy enclosures**
- multichip modules**
- package on package or package in package**
- PCB connections**
- physical protection**
- pin grid arrays**
- postbonding and preseat inspection**
- premade packages**
- quad packages**
- sealing techniques**
- stacking die techniques**
- tape automated bonding (TAB)**

- thin packages
- three-dimensional enabling technologies
- three-dimensional packages
- underfillment
- wire bonding

Pad mask

Pad oxide

PAL (programmable array logic) circuits

Pancake induction-heated APCVD

Parabolic stage of oxidation

Parallel downflow rinser

Parallel processing

Parametric testing

Particles

- as contaminants
- photoresists
- removal of, photomasking and

Particles per wafer pass (PWP)

Particulates

- equipment-induced
- in a gas
- removal of

Passivation layer

Pass-throughs, double-door

Pattern generator

Pattern shift

Patterning (photomasking)

- alignment and exposure
- dark-field masks
- dehydration baking
- final inspection
- goal of
- light-field masks
- mask making
- overview of
- particle removal
- surface preparation
- ten-step process
- (*See also* Etching; Photoresists)

PCB (printed circuit board)

P-channel transistors

PECVD (plasma-enhanced chemical vapor deposition)

Pellicles

Periodic table of the elements

Perkin Elmer scanning projection aligner

Personnel-generated contamination

PH scale

Phase

Phase contrast microscopes

Phase-shift masks

Phenol-formaldehyde polymer

Phenolic organic strippers

Phosphorus silicate glass (PSG)

Photolithography (*see* Patterning)

Photomasking (*see* Patterning)

Photoresist application (spinning)

Photoresists

additives

adhesion capability

basic chemistry of

chemically amplified

cleanliness of

comparison of positive and negative resists

dry etching

exposure speed, sensitivity, and exposure source

light and heat sensitivity of

light-sensitive and energy-sensitive polymers

optical

particle and contamination levels

performance factors

physical properties of

pinholes

process latitude

removal of

resolution capability

sensitizers

shelf life of

soft bake

solvents

step coverage

storage and control of

stripping

thermal flow

Photosolubilization

Physical protection of chips

Physical vapor deposition (PVD)

Pile-up

Pin grid arrays

Pinch resistors

Pinholes

- Piranha etch
- Pitch of the bonding
- Planar plasma etching
- Planar technology
- Planarization
 - etchback
 - global
 - polyimide planarization layers
- Plasma damage
- Plasma descum
- Plasma etching
 - contamination, residues, corrosion, and cost of ownership
- Plasma ion immersion
- Plasma state
- Plasma-enhanced chemical vapor deposition (PECVD)
- Plasma-etch process
- Plastic packages
- Plating process flows
- Plug filling
- Plugs
- PMMA (positive-acting polymethylmethacrylate) resist
- PNP transistors
- Point defects
- Point of use chemical generation (POUCG)
- Point-of-use (POU) chemical mixers
- Poisson model
- Polarization, ellipsometers and
- Polishing:
 - chemical mechanical (*see* Chemical-mechanical polishing)
 - double-sided
 - edge grinding and
 - rough
- Polishing pads, chemical-mechanical polishing (CMP)
- Polishing rates, chemical-mechanical polishing (CMP)
- Polycide-gate MOS
- Polyimide planarization layers
- Polymer, sidewall
- Polymerization
- Polysilicon
 - amorphous silicon deposition and
 - doped
 - oxidation of
- Population
- Portable conformal layer
- Positive and negative resists, comparison of

- Positive resist development
- Positive-acting polymethylmethacrylate (PMMA) resist
- Post-CMP clean
- Post-exposure bake
- Post-ion implant, stripping
- Postoxidation evaluation
- Post-soft-bake cooling (chill)
- POUCG (point of use chemical generation)
- Prebake (hard bake)
- Predep (*see* Deposition)
- Predeposition (*see* Deposition)
- Pre-etch bake (hard bake)
- Premade packages
- Preoxidation wafer cleaning
- Pressure
 - plasma system
- Prime wafers
- Priming
 - spin
 - vacuum vapor
 - vapor
- Printed circuit board (PCB)
- Probe card
- Process automation
- Process chemicals
- Process control techniques
- Process cycle time
- Process defects
- Process improvements
- Process island
- Process latitude
- Process steps, number of
- Process test wafers
- Process variations:
 - wafer diameter and
 - wafer fabrication and
- Process water
- Productivity, yield and
- Profileometers
 - optical
 - surface
- Programmable array logic (PAL) circuits
- Projected range
- Projection aligners
- Projection exposure

Projection systems
PROM (programmable read-only memory)
Properties of matter
Protection of chips
Protons
Proximity aligners
Proximity effect
PSG (phosphorus silicate glass)
P-type dopants
P-type semiconductors
Puddle development
Puddle procedure
Pumps:
 dry mechanical
 turbomolecular hi-vac
PVD (physical vapor deposition)
PWP (particles per wafer pass)
Pyrolysis
Pyrox

Q

Quad packages
Quality control and certification (ISO 9000)
Quartz, as contaminant

R

Radiation
 synchrotron
Radiation damage
Radiation hardened devices
Rail-guided vehicle (RGV)
RAM (random-access memory)
Ramping (temperature ramping)
Rapid thermal oxidation (RTO)
Rapid thermal processing (RTP)
 for post-implant annealing
Raster scanning
Rayleigh constant
Rayleigh formula
RC constant
RCA cleans
Reactive ion etching (RIE)
Read and write capability

- Real-time response
- Recipe
- Reclaim wafers
- Reduction
- Reduction steppers
- Redundancy
- Reflectivity, subsurface
- Reflectometry
- Reflow
- Refraction, index of
- Refractory metals and refractory metal silicides
- Refresh
- Registration
- Registration capability, of aligners
- Reoxidation (reox) (*see* Drive-in oxidation)
- Residues:
 - inorganic
 - organic
- Resist light scattering
- Resist stripping
- Resistance
- Resistivity
 - of doped semiconductors
 - measurements of
- Resistivity check
- Resistors
 - doped
 - EPI
 - ion implantation
 - measurements
 - pinch
 - thin-film
- Resists (*see* Photoresists)
- Resolution
 - of aligners
 - amplified resist and
 - challenges and solutions
 - contrast effects
 - of photoresists
 - scanning electron microscopes
- Resputter
- Reticles
- Retrograde wells
- Reverse sputter
- Reworks (redos)

RGV (rail-guided vehicle)

Rinsing:

after development

dump

sonic-assisted

spray

water

ROM (read-only memory)

Rough polish

Roughing pumps

Run-out and run-in problems

S

Safety issues

Salicide-gate MOS

Same-type doping

Sawing, as die-separation method

Saws

Scaling (die shrink)

SCALPEL

Scan speed

Scanning Auger microanalysis (SAM)

Scanning capacitance microscopy (SCM)

Scanning electron microscope (SEM)

Scanning projection aligners

Scattrometry

Schottky barrier bipolar transistors

Schottky barrier diodes

Scratching

Scribing

Scrubbers, wafer

Scrubbing, in eutectic die attach

Scumming

Seal ring

Sealing techniques

Secondary ion mass spectrometry (SIMS)

Seed deposition

Seeds model

Selection, ion implantation system

Selectivity:

etching

plasma-etching processes

Self-aligned gates

Self-aligned structures

Semiconducting compounds

Semiconductor devices and integrated circuits

capacitors

conductors

diodes

resistors

transistors

(*See also* Integrated circuits)

Semiconductor Equipment and Materials International (SEMI), standards program

Semiconductor industry

advances in fabrication processes

birth of

cost of chips

increase in interconnection levels

increasing chip and wafer sizes

International Technology Roadmap for Semiconductors

invention of integrated circuits (ICs)

Moore's law

nano era

process and product trends

reduction in defect density

solid-state era

stages of manufacturing

Semiconductor Industry Association (SIA)

(*See also International Technology Roadmap for Semiconductors*)

Semiconductor materials

atomic structure

Semiconductors:

deposited

diamond

doped

intrinsic

N-and P-type

production materials

stages of production

Sensitivity of photoresists

Sensitizers, photoresist

Service bays

Shallow trench isolation

Shape metrology (3D shape metrology)

Sheet resistance

Shift

Shockley, William

Shockley, William

Shockley Laboratories

Shoe coverings

Shoe covers

SIA Roadmap (IRTS)

Side diffusion (lateral diffusion)

Sidewall polymer

Signetics

Silane, source chemistry

Silicon

- extraction and purification of strained**

Silicon dioxide

- as device dielectric (MOS gates)**

- doped**

- doping barrier**

- epitaxial**

- as insulator or dielectric**

- surface dielectric**

- surface passivation and**

- thermal growth of**

 - (See also Thermal oxidation)*

- thickness color chart**

- thicknesses**

- wet etching**

Silicon gate MOS transistors

Silicon germanium (SiGe)

Silicon nitride films

Silicon nitride wet etching

Silicon on diamond (SOD)

Silicon on insulator (SOI)

Silicon oxynitride films (SiON; nitrided-oxide or nitrooxide films)

Silicon tetrachloride (SiCl₄)

- chemical vapor deposition of silicon from**

- epitaxial deposition process**

- source chemistry**

Silicon wafers (*see* Wafers)

Silicon wet etching

Silicon-gate MOS

Silox

Silylation (DESIRE process)

SIMOX

Single-layer metal systems

SiP (System in a Package)

Skin flaking

Slip

Slugs
Slurries
Slurry
Small outline IC (SOIC)
Small-scale integration (SSI)
SMIF (standard mechanical interface)
Smog, contamination control and
Snow cleaning
Snowballs
SOC (system-on-a-chip)
Sodium, as mobile ionic contaminant
Soft bake
Soft-contact machines
SOI (silicon on insulator)
Solar cells
Solid neighbor source
Solids, defined
Solids content of photoresists
Solid-state diffusion
Solid-state era
Solutions
Solvent-amine strippers
Solvents
Sonic-assisted cleaning and rinsing
SOS (silicon on sapphire)
Source cabinet
Space charge forces
SPC (statistical process control)
Specific gravity
Spectral response characteristic of resists
Spectrometry
 time of flight secondary ion mass (TOF-SIMS)
Spectroscopes:
 Auger electron
 electron, for chemical analysis (ESCA)
Spikes
Spiking
Spin priming
Spinners
 automatic
 manual
Spinning (photoresist application)
Spin-on-glass (SOG) layer
Spin-rinse dryers
Sporck, Charles

- Spot defects
- Spray cleaning
- Spray development
- Spray etching
- Spray rinsing
- Spreading resistance probe
- Sputter deposit intermetal stack
- Sputter etching
- Sputtering (sputter deposition)
 - argon
 - magnetron
- SRAMs (static random-access memories)
- Stack thickness and composition, evaluation of
- Stacked capacitors
- Stacking die techniques
- Stacking faults
- Stacks, plasma-etching process
- Stained pads
- Staining, electrolytic
- Standard clean-1 (SC-1)
- Standard clean-2 (SC-2)
- Standard logic circuits
- Standards, equipment
- Standing waves
- States of matter
- Static control
 - packaging and
- Static dispense spin process
- Station yields
- Statistical process control (SPC)
- Step and scan aligners
- Step coverage
- Step height measurement
- Steppers
- Stockers
- Stoichiometry
- Storage node, capacitor
- Strain gauges
- Strained silicon
- Stress-relieve oxide (SRO) layer
- Stripping:
 - dry
 - new challenges
 - phenolic organic strippers
 - plasma etch

- post-ion implant
- solvent-amine strippers
- wet chemical

Strontium titanate

Studs

Stylus (surface profileometers)

Subatomic particles

Subcollectors

Subject contrast

Substrate reduction

Subsurface reflectivity

Suckback (drawback)

Sulfolane

Sulfonic acid

Sulfuric acid

- with hydrogen peroxide

- wet chemical stripping

Superconductors

Surface concentration

Surface imaging

Surface inspection, postoxidation

Surface passivation

Surface preparation, photomasking

Surface problems

Surface profileometers (stylus)

Surface roughness

Surface tension drying

Surface tension of photoresists

Surface topography

Surface-mount technology (SMT)

SWAMI process

Switching

Synchrotron radiation

System in a Package (SiP)

System-on-a-chip (SOC)

T

Tape automated bonding (TAB)

Target chamber, end station

Targets (alignment marks)

TC bonding (thermocompression bonding)

TCA (trichloroethane)

TCE (trichloroethylene)

Technology Roadmap (ITRS), packaging technology

TEM (transmission electron microscope)

Temperature

contamination control and

hard bake

room-temperature chemistries

Temperature control system, horizontal tube furnaces

Temperature cycling

Temperature ramping

Temperature sensing

Test head

Test wafers (monitor wafers)

Tests and measurements

Auger electron spectroscopy

burn-in tests

capacitance-voltage profiling

collimated light inspection

concentration or depth profile

contamination and defect detection

critical dimensions and line-width measurements

dark-field inspection

device electrical measurements

device failure analysis-emission microscopy

electrical testing

electron spectroscope for chemical analysis

ellipsometers

final testing

four-point probe

gate oxide integrity electrical measurement

general surface characterization

groove (or bevel) and stain technique

junction depth

layer thickness measurements

microscope techniques

(See also Microscopes)

optical image-shearing dimension measurement

optically modulated optical reflection (thermawave)

physical measurement methods

process and device evaluation

resistivity

scanning capacitance microscopy (SCM)

scanning electron microscope (SEM)

scattrometry

secondary ion mass spectrometry (SIMS)

sheet resistance

spectrophotometers or reflectometry

- spreading resistance probe
- stylus (surface profileometers)
- time of flight secondary ion mass spectrometry (TOF-SIMS)
- ultraviolet light
- visual surface inspection techniques

Tetraethyl orthosilicate (TEOS)

Tetramethylammonium hydroxide (TMAH)

Texas Instruments

Thermal budget

- rapid thermal process (RTP) technology and

Thermal diffusion (*see* Diffusion doping)

Thermal nitridation

Thermal oxidation

- differential oxidation rates and oxide steps
- high-pressure oxidation
- influences on the oxidation rate
- mechanisms of
- methods
- oxidant sources
- oxide impurities
- postoxidation evaluation
- processes
- rapid systems (RTO)
- wafer dopant redistribution and
- wafer orientation and

Thermawave (optically modulated optical reflection)

Thermocompression bonding (TC bonding)

Thermode

Thermosonic gold ball bonding

Thickness:

- film
- stack

Thickness measurements

- ellipsometers and
- scanning electron microscope (SEM)
- (*See also* Tests and measurements)

Thin packages

Thin small outline packages (TSOPs)

Thin-film resistors

Thinning, wafer

Third optical inspection

Three-dimensional enabling technologies

Three-dimensional packages

Threshold voltage (inversion voltage)

- ion implantation

Through-silicon vias (TSVs)
Time of flight secondary ion mass spectrometry (TOF-SIMS)
Tin-lead solder
Titanium nitride (TiN) layers
Titanium-tungsten (TiW) layers
TMAH (tetramethylammonium hydroxide)
TOF-SIMS (time of flight secondary ion mass spectrometry)
Top surface imaged (TSI) techniques
Torr
Tracks, automatic spinners and
Transfer resistor
Transistors
 bipolar
 field-effect
 NPN
 operational analogy
 PNP
Transmission electron microscope (TEM)
Transport-limited reaction
Trench capacitors
Trench isolation
Triboelectric charging
Trichloroethane (TCA)
Trichloroethylene (TCE)
 negative resists and
Trichlorosilane
Tri-gate transistor
Trimming leads
TSVs (through-silicon vias)
TTL (through the lens) system
Tube furnaces:
 horizontal
 vertical
Tungsten, CVD refractory deposition
Tunnel design
Tunneling
Turbomolecular hi-vac pumps
Twining

U

UBM (under-bump-metallization) stack
UHV/CVD (ultrahigh vacuum/chemical vapor deposition)
ULPA (ultra-low-particle) filters
Ultra large-scale integration (ULSI)

Ultrahigh vacuum CVD (UHV-CVD)
Ultra-shallow junctions
Ultrasonic bonding
Ultrasonic cleaning
Ultrathin body (UTB) for MOSFET devices
Ultrathin MOSFET gate thickness
Ultrathin packages (UTPs)
Ultraviolet light inspection
Ultraviolet (UV) ozone
Under-bump-metallization (UBM) stack
Undercutting
Underfillment
Underpass conductors
Unit cells
Unyielded die cost



Vacancies
Vacancy crystal defect
Vacuum baking
Vacuum cleaners
Vacuum evaporation
Vacuum systems
Vacuum tubes
Vacuum vapor priming
Van de Graff, Robert
Van der Pauw structure
Van der Waals force
Vapor, water, control of
Vapor density
Vapor drying
Vapor etching
Vapor-or gas-phase cleaning
Vapor priming
Vapor-phase epitaxy (VPE)
Vapox
Variable numerical aperture lenses
Vector scanning
Vertical laminar flow (VLF) stations
Vertical tube furnaces
Very very large-scale integration (VVLSI)
Vias
Viewing angle
Viscosity of photoresists

Visual surface inspection techniques

VLF (vertical laminar flow) stations

Volatile memories

Voltage:

diode

threshold (inversion voltage)

VPE (vapor-phase epitaxy)

W

Wafer charging

Wafer fabrication

automation

basic operations

batch versus single-wafer processing

book-to-bill ratio (b/b)

breakage and warping

cost of ownership

costs of

die separation

equipment

equipment standards

fab floor layout

goal of

green fabs

inventory control

labor costs

line organization

materials

Moore's law and

prebonding wafer preparation

process variation

production cost factors

quality control and certification (ISO 9000)

statistical process control

steps in

terminology

wafer thinning

wafer-delivery automation

yield and productivity

yield improvements

yield limiters

yields

(See also Packaging; Yields)

Wafer holders

Wafer isolation

Wafer orientation, oxidation growth rate and

Wafer preparation

Wafer probers

Wafer scale packaging (WSP)

Wafer scrubbers

Wafer sort

yields and

Wafer thinning

Wafer throughput

Wafer warping

Wafer-delivery automation

Wafer-loading automation

Wafer-on-wafer technology

Wafers:

automatic wafer loading

backside processing

cleaning (*see* Cleaning)

diameter of (*see* Diameter of wafers)

dislocations

dopant redistribution

drying

evaluation of

increasing diameters of

manual wafer handling

marking

orientation indicators

oxidation (*see* Oxidation)

packaging (*see* Packaging)

polishing (*see* Polishing)

preparation of

prime

priming

quality

slicing

tests and measurements of (*see* Tests and measurements)

(*See also specific topics*)

Wafers in process (WIP)

Wafer-surface cleaning

Warping, wafer

Water, process

Water rinsing

Water spray

high-pressure

Water vapor:

**control of
sources of**

Watermarks

Watson, Thomas

Wedge bonding

Weight

Wet chemical stripping

Wet development processes

Wet etching

aluminum-film

deposited-oxide

silicon

silicon dioxide

silicon nitride

vapor etching

Wet processes (wet cleaning)

Wet-spray etching

Wire bonding

aluminum wire

gold wire

sputter deposit intermetal stack

Wire saws

Wiring

WIT (wafer isolation technology)

Work cells

Work function

Workstations, cleanroom techniques

WSP (near wafer scale packaging)

X

X-R control chart

X-ray blocking masks

X-ray resists

X-rays, as exposure sources

Y

Yield enhancement

Yield measurement points

Yield models

Yielded die cost

Yields

accumulative wafer-fabrication

assembly and final test

circuit density and defect density
die area and defect density
feature size and defect size
improvements in
mask defects and
number of process steps
overall process
process cycle time
process defects and
process variation and
productivity and
station
wafer breakage and warping
wafer diameter and crystal defects
wafer diameter and die size
wafer diameter and edge die
wafer diameter and process variations
wafer-fabrication yield limiters
wafer-sort yield factors
wafer-sort yield formulas
(*See also* Wafer fabrication)

Z

Zeta potential